

UNIVERSITY OF TORONTO



3 1761 00184469 5



Digitized by the Internet Archive
in 2007 with funding from
Microsoft Corporation



ENGINEERING MATHEMATICS

McGraw-Hill Book Co. Inc.

PUBLISHERS OF BOOKS FOR

Coal Age ▽ Electric Railway Journal
Electrical World ▽ Engineering News-Record
American Machinist ▽ Ingeniería Internacional
Engineering & Mining Journal ▽ Power
Chemical & Metallurgical Engineering
Electrical Merchandising

ENGINEERING MATHEMATICS

A SERIES OF LECTURES DELIVERED
AT UNION COLLEGE

BY

CHARLES PROTEUS STEINMETZ, A.M., PH.D.

PAST PRESIDENT

AMERICAN INSTITUTE OF ELECTRICAL ENGINEERS

THIRD EDITION
REVISED AND ENLARGED
SECOND IMPRESSION

172188
20/6/22

McGRAW-HILL BOOK COMPANY, INC.
239 WEST 39TH STREET. NEW YORK

LONDON: HILL PUBLISHING CO., LTD.
6 & 8 BOUVERIE ST., E. C.

1917

COPYRIGHT, 1911, 1915 AND 1917, BY THE
MCGRAW-HILL BOOK COMPANY, INC.

QA
37
508
1917

PREFACE TO THIRD EDITION.

In preparing the third edition of Engineering Mathematics, besides revision and correction of the previous text, considerable new matter has been added.

The chain fraction has been recognized and discussed as a convenient method of numerical representation and approximation; a paragraph has been devoted to the diophantic equations, and a section added on engineering reports, discussing the different purposes for which engineering reports are made, and the corresponding character and nature of the report, in its bearing on the success and recognition of the engineer's work.

CHARLES PROTEUS STEINMETZ.

CAMP MOHAWK,
September 1st, 1917.



PREFACE TO FIRST EDITION.

THE following work embodies the subject-matter of a lecture course which I have given to the junior and senior electrical engineering students of Union University for a number of years.

It is generally conceded that a fair knowledge of mathematics is necessary to the engineer, and especially the electrical engineer. For the latter, however, some branches of mathematics are of fundamental importance, as the algebra of the general number, the exponential and trigonometric series, etc., which are seldom adequately treated, and often not taught at all in the usual text-books of mathematics, or in the college course of analytic geometry and calculus given to the engineering students, and, therefore, electrical engineers often possess little knowledge of these subjects. As the result, an electrical engineer, even if he possess a fair knowledge of mathematics, may often find difficulty in dealing with problems, through lack of familiarity with these branches of mathematics, which have become of importance in electrical engineering, and may also find difficulty in looking up information on these subjects.

In the same way the college student, when beginning the study of electrical engineering theory, after completing his general course of mathematics, frequently finds himself sadly deficient in the knowledge of mathematical subjects, of which a complete familiarity is required for effective understanding of electrical engineering theory. It was this experience which led me some years ago to start the course of lectures which is reproduced in the following pages. I have thus attempted to bring together and discuss explicitly, with numerous practical applications, all those branches of mathematics which are of special importance to the electrical engineer. Added thereto

are a number of subjects which experience has shown me to be important for the effective and expeditious execution of electrical engineering calculations. Mere theoretical knowledge of mathematics is not sufficient for the engineer, but it must be accompanied by ability to apply it and derive results—to carry out numerical calculations. It is not sufficient to know how a phenomenon occurs, and how it may be calculated, but very often there is a wide gap between this knowledge and the ability to carry out the calculation; indeed, frequently an attempt to apply the theoretical knowledge to derive numerical results leads, even in simple problems, to apparently hopeless complication and almost endless calculation, so that all hope of getting reliable results vanishes. Thus considerable space has been devoted to the discussion of methods of calculation, the use of curves and their evaluation, and other kindred subjects requisite for effective engineering work.

Thus the following work is not intended as a complete course in mathematics, but as supplementary to the general college course of mathematics, or to the general knowledge of mathematics which every engineer and really every educated man should possess.

In illustrating the mathematical discussion, practical examples, usually taken from the field of electrical engineering, have been given and discussed. These are sufficiently numerous that any example dealing with a phenomenon with which the reader is not yet familiar may be omitted and taken up at a later time.

As appendix is given a descriptive outline of the introduction to the theory of functions, since the electrical engineer should be familiar with the general relations between the different functions which he meets.

In relation to "Theoretical Elements of Electrical Engineering," "Theory and Calculation of Alternating Current Phenomena," and "Theory and Calculation of Transient Electric Phenomena," the following work is intended as an introduction and explanation of the mathematical side, and the most efficient method of study, appears to me, to start with "Electrical Engineering Mathematics," and after entering its third chapter, to take up the reading of the first section of "Theoretical Elements," and then parallel the study of "Electrical

Engineering Mathematics," "Theoretical Elements of Electrical Engineering," and "Theory and Calculation of Alternating Current Phenomena," together with selected chapters from "Theory and Calculation of Transient Electric Phenomena," and after this, once more systematically go through all four books.

CHARLES P. STEINMETZ.

SCHENECTADY, N. Y.,
December, 1910.

PREFACE TO SECOND EDITION.

IN preparing the second edition of Engineering Mathematics, besides revision and correction of the previous text, considerable new matter has been added, more particularly with regard to periodic curves. In the former edition the study of the wave shapes produced by various harmonics, and the recognition of the harmonics from the wave shape, have not been treated, since a short discussion of wave shapes is given in "Alternating Current Phenomena." Since, however, the periodic functions are the most important in electrical engineering, it appears necessary to consider their shape more extensively, and this has been done in the new edition.

The symbolism of the general number, as applied to alternating waves, has been changed in conformity to the decision of the International Electrical Congress of Turin, a discussion of the logarithmic and semi-logarithmic scale of curve plotting given, etc.

CHARLES P. STEINMETZ.

December, 1914.



CONTENTS.

PREFACE	PAGE V
---------------	-----------

CHAPTER I. THE GENERAL NUMBER.

A. THE SYSTEM OF NUMBERS.

1. Addition and Subtraction. Origin of numbers. Counting and measuring. Addition. Subtraction as reverse operation of addition.....	1
2. Limitation of subtraction. Subdivision of the absolute numbers into positive and negative.	2
3. Negative number a mathematical conception like the imaginary number. Cases where the negative number has a physical meaning, and cases where it has not	4
4. Multiplication and Division. Multiplication as multiple addition. Division as its reverse operation. Limitation of division.....	6
5. The fraction as mathematical conception. Cases where it has a physical meaning, and cases where it has not	8
6. Involution and Evolution. Involution as multiple multiplication. Evolution as its reverse operation. Negative exponents.	9
7. Multiple involution leads to no new operation.....	10
8. Fractional exponents.....	10
9. Irrational Numbers. Limitation of evolution. Endless decimal fraction. Rationality of the irrational number. . .	11
10. Quadrature numbers. Multiple values of roots. Square root of negative quantity representing quadrature number, or rotation by 90° . .	13
11. Comparison of positive, negative and quadrature numbers. Reality of quadrature number. Cases where it has a physical meaning, and cases where it has not	14
12. General Numbers. Representation of the plane by the general number. Its relation to rectangular coordinates. . . .	16
13. Limitation of algebra by the general number. Roots of the unit. Number of such roots, and their relation	18
14. The two reverse operations of involution	19

	PAGE
15. Logarithmation. Relation between logarithm and exponent of involution. Reduction to other base. Logarithm of negative quantity	20
16. Quaternions. Vector calculus of space.....	22
17. Space rotors and their relation. Super algebraic nature of space analysis.....	22
B. ALGEBRA OF THE GENERAL NUMBER OF COMPLEX QUANTITY.	
Rectangular and Polar Coordinates.....	25
18. Powers of j . Ordinary or real, and quadrature or imaginary number. Relations	25
19. Conception of general number by point of plane in rectangular coordinates; in polar coordinates. Relation between rectangular and polar form.....	26
20. Addition and Subtraction. Algebraic and geometrical addition and subtraction. Combination and resolution by parallelogram law	28
21. Denotations	30
22. Sign of vector angle. Conjugate and associate numbers. Vector analysis	30
23. Instance of steam path of turbine.....	33
24. Multiplication. Multiplication in rectangular coordinates....	38
25. Multiplication in polar coordinates. Vector and operator	38
26. Physical meaning of result of algebraic operation. Representation of result	40
27. Limitation of application of algebraic operations to physical quantities, and of the graphical representation of the result. Graphical representation of algebraic operations between current, voltage and impedance	40
28. Representation of vectors and of operators	42
29. Division. Division in rectangular coordinates	42
30. Division in polar coordinates	43
31. Involution and Evolution. Use of polar coordinates	44
32. Multiple values of the result of evolution. Their location in the plane of the general number. Polyphase and n phase systems of numbers	45
33. The n values of $\sqrt[n]{1}$ and their relation.....	46
34. Evolution in rectangular coordinates. Complexity of result ...	47
35. Reduction of products and fractions of general numbers by polar representation. Instance.....	48
36. Exponential representations of general numbers. The different forms of the general number	49
37. Instance of use of exponential form in solution of differential equation	50

	PAGE
38. Logarithmation. Resolution of the logarithm of a general number	51

CHAPTER II. THE POTENTIAL SERIES AND EXPONENTIAL FUNCTION.

A. GENERAL.

39. The infinite series of powers of x	52
40. Approximation by series.....	53
41. Alternate and one-sided approximation	54
42. Convergent and divergent series.....	55
43. Range of convergency. Several series of different ranges for same expression	56
44. Discussion of convergency in engineering applications	57
45. Use of series for approximation of small terms. Instance of electric circuit	58
46. Binomial theorem for development in series. Instance of inductive circuit.....	59
47. Necessity of development in series. Instance of arc of hyperbola	60
48. Instance of numerical calculation of $\log (1+x)$	63

B. DIFFERENTIAL EQUATIONS.

49. Character of most differential equations of electrical engineering. Their typical forms	64
50. $\frac{dy}{dx} = y$. Solution by series, by method of indeterminate coefficients.....	65
51. $\frac{dz}{dx} = az$. Solution by indeterminate coefficients.....	68
52. Integration constant and terminal conditions	68
53. Involution of solution. Exponential function	70
54. Instance of rise of field current in direct current shunt motor ..	72
55. Evaluation of inductance, and numerical calculation	75
56. Instance of condenser discharge through resistance	76
57. Solution of $\frac{d^2y}{dx^2} = ay$ by indeterminate coefficients, by exponential function.....	78
58. Solution by trigonometric functions	81
59. Relations between trigonometric functions and exponential functions with imaginary exponent, and inversely	83
60. Instance of condenser discharge through inductance. The two integration constants and terminal conditions	84
61. Effect of resistance on the discharge. The general differential equation.....	86

	PAGE
62. Solution of the general differential equation by means of the exponential function, by the method of indeterminate coefficients.....	86
63. Instance of condenser discharge through resistance and inductance. Exponential solution and evaluation of constants....	88
64. Imaginary exponents of exponential functions. Reduction to trigonometric functions. The oscillating functions.....	91
65. Explanation of tables of exponential functions).....	92

CHAPTER III. TRIGONOMETRIC SERIES.

A. TRIGONOMETRIC FUNCTIONS.

66. Definition of trigonometric functions on circle and right triangle	94
67. Sign of functions in different quadrants	95
68. Relations between sin, cos, tan and cot	97
69. Negative, supplementary and complementary angles	98
70. Angles $(x \pm \pi)$ and $\left(x \pm \frac{\pi}{2}\right)$	100
71. Relations between two angles, and between angle and double angle	102
72. Differentiation and integration of trigonometric functions. Definite integrals.....	103
73. The binomial relations	104
74. Polyphase relations.....	104
75. Trigonometric formulas of the triangle	105

B. TRIGONOMETRIC SERIES.

76. Constant, transient and periodic phenomena. Univalent periodic function represented by trigonometric series.....	106
77. Alternating sine waves and distorted waves	107
78. Evaluation of the Constants from Instantaneous Values. Calculation of constant term of series	108
79. Calculation of cos-coefficients	110
80. Calculation of sin-coefficients	113
81. Instance of calculating 11th harmonic of generator wave	114
82. Discussion. Instance of complete calculation of pulsating current wave	116
83. Alternating waves as symmetrical waves. Calculation of symmetrical wave	117
84. Separation of odd and even harmonics and of constant term ...	120
85. Separation of sine and cosine components	121
86. Separation of wave into constant term and 4 component waves	122
87. Discussion of calculation	123
88. Mechanism of calculation.....	124

CONTENTS.

xv

	PAGE
89. Instance of resolution of the annual temperature curve	125
90. Constants and equation of temperature wave	131
91. Discussion of temperature wave	132
C. REDUCTION OF TRIGONOMETRIC SERIES BY POLYPHASE RELATION.	
92. Method of separating certain classes of harmonics, and its limitation	134
93. Instance of separating the 3d and 9th harmonic of transformer exciting current	136
D. CALCULATION OF TRIGONOMETRIC SERIES FROM OTHER TRIGONOMETRIC SERIES.	
94. Instance of calculating current in long distance transmission line, due to distorted voltage wave of generator. Line constants. .	139
95. Circuit equations, and calculation of equation of current	141
96. Effective value of current, and comparison with the current produced by sine wave of voltage	143
97. Voltage wave of reactance in circuit of this distorted current ...	145

CHAPTER IV. MAXIMA AND MINIMA.

98. Maxima and minima by curve plotting. Instance of magnetic permeability. Maximum power factor of induction motor as function of load	147
99. Interpolation of maximum value in method of curve plotting. Error in case of unsymmetrical curve. Instance of efficiency of steam turbine nozzle. Discussion	149
100. Mathematical method. Maximum, minimum and inflexion point. Discussion	152
101. Instance: Speed of impulse turbine wheel for maximum efficiency. Current in transformer for maximum efficiency. .	154
102. Effect of intermediate variables. Instance: Maximum power in resistance shunting a constant resistance in a constant current circuit	155
103. Simplification of calculation by suppression of unnecessary terms, etc. Instance	157
104. Instance: Maximum non-inductive load on inductive transmission line. Maximum current in line	158
105. Discussion of physical meaning of mathematical extremum. Instance	160
106. Instance: External reactance giving maximum output of alternator at constant external resistance and constant excitation. Discussion	161
107. Maximum efficiency of alternator on non-inductive load. Discussion of physical limitations	162

	PAGE
108. Extrema with several independent variables. Method of mathematical calculation, and geometrical meaning	163
109. Resistance and reactance of load to give maximum output of transmission line, at constant supply voltage	165
110. Discussion of physical limitations	167
111. Determination of extrema by plotting curve of differential quotient. Instance: Maxima of current wave of alternator of distorted voltage on transmission line	168
112. Graphical calculation of differential curve of empirical curve, for determining extrema	170
113. Instance: Maximum permeability calculation	170
114. Grouping of battery cells for maximum power in constant resistance	171
115. Voltage of transformer to give maximum output at constant loss	173
116. Voltage of transformer, at constant output, to give maximum efficiency at full load, at half load.	174
117. Maximum value of charging current of condenser through inductive circuit (a) at low resistance; (b) at high resistance.	175
118. At what output is the efficiency of an induction generator a maximum?	177
119. Discussion of physical limitations. Maximum efficiency at constant current output.	178
120. METHOD OF LEAST SQUARES. Relation of number of observations to number of constants. Discussion of errors of observation.	179
121. Probability calculus and the minimum sum of squares of the errors.	181
122. The differential equations of the sum of least squares	182
123. Instance: Reduction of curve of power of induction motor running light, into the component losses. Discussion of results	182
123A. Diophantic equations.	186

CHAPTER V. METHODS OF APPROXIMATION.

124. Frequency of small quantities in electrical engineering problems. Instances. Approximation by dropping terms of higher order.	187
125. Multiplication of terms with small quantities	188
126. Instance of calculation of power of direct current shunt motor	189
127. Small quantities in denominator of fractions	190
128. Instance of calculation of induction motor current, as function of slip	191

	PAGE
129. Use of binomial series in approximations of powers and roots, and in numerical calculations	193
130. Instance of calculation of current in alternating circuit of low inductance. Instance of calculation of short circuit current of alternator, as function of speed	195
131. Use of exponential series and logarithmic series in approximations	196
132. Approximations of trigonometric functions	198
133. McLaurin's and Taylor's series in approximations	198
134. Tabulation of various infinite series and of the approximations derived from them	199
135. Estimation of accuracy of approximation. Application to short circuit current of alternator	200
136. Expressions which are approximated by $(1+s)$ and by $(1-s)$	201
137. Mathematical instance of approximation	203
138. EQUATIONS OF THE TRANSMISSION LINE. Integration of the differential equations.	204
139. Substitution of the terminal conditions	205
140. The approximate equations of the transmission line	206
141. Numerical instance. Discussion of accuracy of approximation.	207
141A. Approximation by chain fraction.	208
141B. Approximation by chain fraction	208c

CHAPTER VI. EMPIRICAL CURVES.

A. GENERAL.

142. Relation between empirical curves, empirical equations and rational equations	209
143. Physical nature of phenomenon. Points at zero and at infinity. Periodic or non-periodic. Constant terms. Change of curve law. Scale.	210

B. NON-PERIODIC CURVES.

144. Potential Series. Instance of core-loss curve	212
145. Rational and irrational use of potential series. Instance of fan motor torque. Limitations of potential series	214
146. PARABOLIC AND HYPERBOLIC CURVES. Various shapes of parabolas and of hyperbolas.	216
147. The characteristic of parabolic and hyperbolic curves. Its use and limitation by constant terms.	223
148. The logarithmic characteristic. Its use and limitation	224
149. EXPONENTIAL AND LOGARITHMIC CURVES. The exponential function	227
150. Characteristics of the exponential curve, their use and limitation by constant term. Comparison of exponential curve and hyperbola	228

	PAGE
151. Double exponential functions. Various shapes thereof.....	231
152. EVALUATION OF EMPIRICAL CURVES. General principles of investigation of empirical curves.....	233
153. Instance: The volt-ampere characteristic of the tungsten lamp, reduced to parabola with exponent 0.6. Rationalized by reduction to radiation law.....	235
154. The volt-ampere characteristic of the magnetite arc, reduced to hyperbola with exponent 0.5.....	238
155. Change of electric current with change of circuit conditions, reduced to double exponential function of time.....	241
156. Rational reduction of core-loss curve of paragraph 144, by parabola with exponent 1.6.	244
157. Reduction of magnetic characteristic, for higher densities, to hyperbolic curve. Instance of the investigation of a hysteresis curve of silicon steel.....	246
C. PERIODIC CURVES.	
158. Distortion of sine wave by harmonics.....	255
159. Third and fifth harmonic. Peak, multiple peak, flat top and sawtooth.....	255
160. Combined effect of third and fifth harmonic.....	263
161. Even harmonics. Unequal shape and length of half waves. Combined second and third harmonic.....	266
162. Effect of high harmonics.....	269
163. Ripples and nodes caused by higher harmonics. Incommensurable waves.....	271

CHAPTER VII. NUMERICAL CALCULATIONS.

164. METHOD OF CALCULATION. Tabular form of calculation.....	275
165. Instance of transmission line regulation.....	277
166. EXACTNESS OF CALCULATION. Degrees of exactness: magnitude, approximate, exact.....	279
167. Number of decimals.....	281
168. INTELLIGIBILITY OF ENGINEERING DATA. Curve plotting for showing shape of function, and for record of numerical values	283
169. Scale of curves. Principles.....	286
170. Logarithmic and semi-logarithmic paper and its proper use.....	287
171. Completeness of record.....	290
171A. Engineering Reports.....	290
172. RELIABILITY OF NUMERICAL CALCULATIONS. Necessity of reliability in engineering calculations.....	293
173. Methods of checking calculations. Curve plotting.....	293a
174. Some frequent errors.....	293b

APPENDIX A. NOTES ON THE THEORY OF FUNCTIONS.

A. GENERAL FUNCTIONS.

175. Implicit analytic function. Explicit analytic function. Reverse function.....	294
--	-----

	PAGE
176. Rational function. Integer function. Approximations by Taylor's Theorem	295
177. Abelian integrals and Abelian functions. Logarithmic integral and exponential functions	296
178. Trigonometric integrals and trigonometric functions. Hyperbolic integrals and hyperbolic functions	297
179. Elliptic integrals and elliptic functions. Their double periodicity	298
180. Theta functions. Hyperelliptic integrals and functions	300
181. Elliptic functions in the motion of the pendulum and the surging of synchronous machines	301
182. Instance of the arc of an ellipsis	301
B. SPECIAL FUNCTIONS.	
183. Infinite summation series. Infinite product series	302
184. Functions by integration. Instance of the propagation functions of electric waves and impulses	303
185. Functions defined by definite integrals	305
186. Instance of the gamma function	306
C. EXPONENTIAL, TRIGONOMETRIC AND HYPERBOLIC FUNCTIONS.	
187. Functions of real variables	306
188. Functions of imaginary variables	308
189. Functions of complex variables	308
190. Relations	309

APPENDIX B. TABLES.

TABLE I. Three decimal exponential functions	312
TABLE II. Logarithms of exponential functions	
Exponential functions	313
Hyperbolic functions	314
INDEX	315



ENGINEERING MATHEMATICS.

CHAPTER I.

THE GENERAL NUMBER.

A. THE SYSTEM OF NUMBERS.

Addition and Subtraction.

1. From the operation of counting and measuring arose the art of figuring, arithmetic, algebra, and finally, more or less, the entire structure of mathematics.

During the development of the human race throughout the ages, which is repeated by every child during the first years of life, the first conceptions of numerical values were vague and crude: many and few, big and little, large and small. Later the ability to count, that is, the knowledge of numbers, developed, and last of all the ability of measuring, and even up to-day, measuring is to a considerable extent done by counting: steps, knots, etc.

From counting arose the simplest arithmetical operation—*addition*. Thus we may count a bunch of horses:

1, 2, 3, 4, 5,

and then count a second bunch of horses,

1, 2, 3;

now put the second bunch together with the first one, into one bunch, and count them. That is, after counting the horses

of the first bunch, we continue to count those of the second bunch, thus:

1, 2, 3, 4, 5,—6, 7, 8;

which gives addition,

$$5 + 3 = 8;$$

or, in general,

$$a + b = c.$$

We may take away again the second bunch of horses, that means, we count the entire bunch of horses, and then count off those we take away thus:

1, 2, 3, 4, 5, 6, 7, 8—7, 6, 5;

which gives *subtraction*,

$$8 - 3 = 5;$$

or, in general,

$$c - b = a.$$

The reverse of putting a group of things together with another group is to take a group away, therefore subtraction is the reverse of addition.

2. Immediately we notice an essential difference between addition and subtraction, which may be illustrated by the following examples:

Addition: 5 horses + 3 horses gives 8 horses,

Subtraction: 5 horses — 3 horses gives 2 horses,

Addition: 5 horses + 7 horses gives 12 horses,

Subtraction: 5 horses — 7 horses is impossible.

From the above it follows that we can always add, but we cannot always subtract; subtraction is not always possible; it is not, when the number of things which we desire to subtract is greater than the number of things from which we desire to subtract.

The same relation obtains in measuring; we may measure a distance from a starting point *A* (Fig. 1), for instance in steps, and then measure a second distance, and get the total distance from the starting point by addition: 5 steps. from *A* to *B*.

then 3 steps, from *B* to *C*, gives the distance from *A* to *C*, as 8 steps.

$$5 \text{ steps} + 3 \text{ steps} = 8 \text{ steps};$$

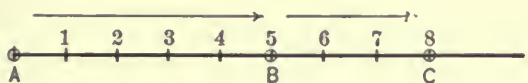


FIG. 1. Addition.

or, we may step off a distance, and then step back, that is, subtract another distance, for instance (Fig. 2),

$$5 \text{ steps} - 3 \text{ steps} = 2 \text{ steps};$$

that is, going 5 steps, from *A* to *B*, and then 3 steps back, from *B* to *C*, brings us to *C*, 2 steps away from *A*.

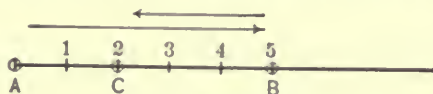


FIG. 2. Subtraction.

Trying the case of subtraction which was impossible, in the example with the horses, $5 \text{ steps} - 7 \text{ steps} = ?$ We go from the starting point, *A*, 5 steps, to *B*, and then step back 7 steps; here we find that sometimes we can do it, sometimes we cannot do it; if back of the starting point *A* is a stone wall, we cannot step back 7 steps. If *A* is a chalk mark in the road, we may step back beyond it, and come to *C* in Fig. 3. In the latter case,

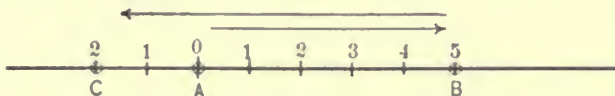


FIG. 3. Subtraction, Negative Result.

at *C* we are again 2 steps distant from the starting point, just as in Fig. 2. That is,

$$5 - 3 = 2 \quad (\text{Fig. 2}),$$

$$5 - 7 = 2 \quad (\text{Fig. 3}).$$

In the case where we can subtract 7 from 5, we get the same distance from the starting point as when we subtract 3 from 5,

but the distance AC in Fig. 3, while the same, 2 steps, as in Fig. 2, is different in character, the one is toward the left, the other toward the right. That means, we have two kinds of distance units, those to the right and those to the left, and have to find some way to distinguish them. The distance 2 in Fig. 3 is toward the left of the starting point A , that is, in that direction, in which we step when subtracting, and it thus appears natural to distinguish it from the distance 2 in Fig. 2, by calling the former -2 , while we call the distance $\overline{AC}$ in Fig. 2: $+2$, since it is in the direction from A , in which we step in adding.

This leads to a subdivision of the system of absolute numbers,

$$1, 2, 3, \dots$$

into two classes, positive numbers,

$$+1, +2, +3, \dots;$$

and negative numbers,

$$-1, -2, -3, \dots;$$

and by the introduction of negative numbers, we can always carry out the mathematical operation of subtraction:

$$c - b = a,$$

and, if b is greater than c , a merely becomes a negative number.

3. We must therefore realize that the negative number and the negative unit, -1 , is a mathematical fiction, and not in universal agreement with experience, as the absolute number found in the operation of counting, and the negative number does not always represent an existing condition in practical experience.

In the application of numbers to the phenomena of nature, we sometimes find conditions where we can give the negative number a physical meaning, expressing a relation as the reverse to the positive number; in other cases we cannot do this. For instance, 5 horses -7 horses = -2 horses has no physical meaning. There exist no negative horses, and at the best we could only express the relation by saying, 5 horses -7 horses is impossible, 2 horses are missing.

In the same way, an illumination of 5 foot-candles, lowered by 3 foot-candles, gives an illumination of 2 foot-candles, thus,

$$5 \text{ foot-candles} - 3 \text{ foot-candles} = 2 \text{ foot-candles.}$$

If it is tried to lower the illumination of 5 foot-candles by 7 foot-candles, it will be found impossible; there cannot be a negative illumination of 2 foot-candles; the limit is zero illumination, or darkness.

From a string of 5 feet length, we can cut off 3 feet, leaving 2 feet, but we cannot cut off 7 feet, leaving -2 feet of string.

In these instances, the negative number is meaningless, a mere imaginary mathematical fiction.

If the temperature is 5 deg. cent. above freezing, and falls 3 deg., it will be 2 deg. cent. above freezing. If it falls 7 deg. it will be 2 deg. cent. below freezing. The one case is just as real physically, as the other, and in this instance we may express the relation thus:

$$+5 \text{ deg. cent.} - 3 \text{ deg. cent.} = +2 \text{ deg. cent.,}$$

$$+5 \text{ deg. cent.} - 7 \text{ deg. cent.} = -2 \text{ deg. cent.;}$$

that is, in temperature measurements by the conventional temperature scale, the negative numbers have just as much physical existence as the positive numbers.

The same is the case with time, we may represent future time, from the present as starting point, by positive numbers, and past time then will be represented by negative numbers. But we may equally well represent past time by positive numbers, and future times then appear as negative numbers. In this, and most other physical applications, the negative number thus appears equivalent with the positive number, and interchangeable: we may choose any direction as positive, and the reverse direction then is negative. Mathematically, however, a difference exists between the positive and the negative number; the positive unit, multiplied by itself, remains a positive unit, but the negative unit, multiplied with itself, does not remain a negative unit, but becomes positive:

$$(+1) \times (+1) = (+1):$$

$$(-1) \times (-1) = (+1), \text{ and not } = (-1).$$

Starting from 5 deg. northern latitude and going 7 deg. south, brings us to 2 deg. southern latitude, which may be expressed thus,

$$+5 \text{ deg. latitude} - 7 \text{ deg. latitude} = -2 \text{ deg. latitude.}$$

Therefore, in all cases, where there are two opposite directions, right and left on a line, north and south latitude, east and west longitude, future and past, assets and liabilities, etc., there may be application of the negative number; in other cases, where there is only one kind or direction, counting horses, measuring illumination, etc., there is no physical meaning which would be represented by a negative number. There are still other cases, where a meaning may sometimes be found and sometimes not; for instance, if we have 5 dollars in our pocket, we cannot take away 7 dollars; if we have 5 dollars in the bank, we may be able to draw out 7 dollars, or we may not, depending on our credit. In the first case, 5 dollars -7 dollars is an impossibility, while the second case 5 dollars -7 dollars = 2 dollars overdraft.

In any case, however, we must realize that the negative number is not a physical, but a mathematical conception, which may find a physical representation, or may not, depending on the physical conditions to which it is applied. The negative number thus is just as imaginary, and just as real, depending on the case to which it is applied, as the imaginary number $\sqrt{-1}$, and the only difference is, that we have become familiar with the negative number at an earlier age, where we were less critical, and thus have taken it for granted, become familiar with it by use, and usually do not realize that it is a mathematical conception, and not a physical reality. When we first learned it, however, it was quite a step to become accustomed to saying, $5-7 = -2$, and not simply, $5-7$ is impossible.

Multiplication and Division.

4. If we have a bunch of 4 horses, and another bunch of 4 horses, and still another bunch of 4 horses, and add together the three bunches of 4 horses each, we get,

$$4 \text{ horses} + 4 \text{ horses} + 4 \text{ horses} = 12 \text{ horses;}$$

or, as we express it,

$$3 \times 4 \text{ horses} = 12 \text{ horses.}$$

The operation of multiple addition thus leads to the next operation, *multiplication*. Multiplication is multiple addition,

$$b \times a = c,$$

thus means

$$a + a + a + \dots (b \text{ terms}) = c.$$

Just like addition, multiplication can always be carried out.

Three groups of 4 horses each, give 12 horses. Inversely, if we have 12 horses, and divide them into 3 equal groups, each group contains 4 horses. This gives us the reverse operation of multiplication, or *division*, which is written, thus:

$$\frac{12 \text{ horses}}{3} = 4 \text{ horses;}$$

or, in general,

$$\frac{c}{b} = a.$$

If we have a bunch of 12 horses, and divide it into two equal groups, we get 6 horses in each group, thus:

$$\frac{12 \text{ horses}}{2} = 6 \text{ horses,}$$

if we divide into 4 equal groups,

$$\frac{12 \text{ horses}}{4} = 3 \text{ horses.}$$

If now we attempt to divide the bunch of 12 horses into 5 equal groups, we find we cannot do it; if we have 2 horses in each group, 2 horses are left over; if we put 3 horses in each group, we do not have enough to make 5 groups; that is, 12 horses divided by 5 is impossible; or, as we usually say; 12 horses divided by 5 gives 2 horses and 2 horses left over, which is written,

$$\frac{12}{5} = 2, \text{ remainder } 2.$$

Thus it is seen that the reverse operation of multiplication, or division, cannot always be carried out.

5. If we have 10 apples, and divide them into 3, we get 3 apples in each group, and one apple left over.

$$\frac{10}{3} = 3, \text{ remainder } 1,$$

we may now cut the left-over apple into 3 equal parts, in which case,

$$\frac{10}{3} = 3 + \frac{1}{3} = 3\frac{1}{3}.$$

In the same manner, if we have 12 apples, we can divide into 5, by cutting 2 apples each into 5 equal pieces, and get in each of the 5 groups, 2 apples and 2 pieces.

$$\frac{12}{5} = 2 + 2 \times \frac{1}{5} = 2\frac{2}{5}.$$

To be able to carry the operation of division through for all numerical values, makes it necessary to introduce a new unit, smaller than the original unit, and derived as a part of it.

Thus, if we divide a string of 10 feet length into 3 equal parts, each part contains 3 feet, and 1 foot is left over. One foot is made up of 12 inches, and 12 inches divided into 3 gives 4 inches; hence, 10 feet divided by 3 gives 3 feet 4 inches.

Division leads us to a new form of numbers: the fraction.

The fraction, however, is just as much a mathematical conception, which sometimes may be applicable, and sometimes not, as the negative number. In the above instance of 12 horses, divided into 5 groups, it is not applicable.

$$\frac{12 \text{ horses}}{5} = 2\frac{2}{5} \text{ horses}$$

is impossible; we cannot have fractions of horses, and what we would get in this attempt would be 5 groups, each comprising 2 horses and some pieces of carcass.

Thus, the mathematical conception of the fraction is applicable to those physical quantities which can be divided into smaller units, but is not applicable to those, which are indivisible, or *individuals*, as we usually call them.

Involution and Evolution.

6. If we have a product of several equal factors, as,

$$4 \times 4 \times 4 = 64,$$

it is written as,

$$4^3 = 64;$$

or, in general,

$$a^b = c.$$

The operation of multiple multiplication of equal factors thus leads to the next algebraic operation—*involution*; just as the operation of multiple addition of equal terms leads to the operation of multiplication.

The operation of involution, defined as multiple multiplication, requires the exponent b to be an integer number; b is the number of factors.

Thus 4^{-3} has no immediate meaning; it would by definition be 4 multiplied (-3) times with itself.

Dividing continuously by 4, we get, $4^6 \div 4 = 4^5$; $4^5 \div 4 = 4^4$; $4^4 \div 4 = 4^3$; etc., and if this successive division by 4 is carried still further, we get the following series:

$$\frac{4^3}{4} = \frac{4 \times 4 \times 4}{4} = 4 \times 4 = 4^2$$

$$\frac{4^2}{4} = \frac{4 \times 4}{4} = 4 = 4^1$$

$$\frac{4^1}{4} = \frac{4}{4} = 1 = 4^0$$

$$\frac{4^0}{4} = \frac{1}{4} = \frac{1}{4} = 4^{-1}$$

$$\frac{4^{-1}}{4} = \frac{1}{4} \div 4 = \frac{1}{4 \times 4} = 4^{-2} = \frac{1}{4^2}$$

$$\frac{4^{-2}}{4} = \frac{1}{4^2} \div 4 = \frac{1}{4 \times 4 \times 4} = 4^{-3} = \frac{1}{4^3};$$

or, in general,

$$a^{-b} = \frac{1}{a^b},$$

$$a^0 = 1.$$

Thus, powers with negative exponents, as a^{-b} , are the reciprocals of the same powers with positive exponents: $\frac{1}{a^b}$.

7. From the definition of involution then follows,

$$a^b \times a^n = a^{b+n},$$

because a^b means the product of b equal factors a , and a^n the product of n equal factors a , and $a^b \times a^n$ thus is a product having $b+n$ equal factors a . For instance,

$$4^3 \times 4^2 = (4 \times 4 \times 4) \times (4 \times 4) = 4^5.$$

The question now arises, whether by multiple involution we can reach any further mathematical operation. For instance,

$$(4^3)^2 = ?,$$

may be written,

$$\begin{aligned} (4^3)^2 &= 4^3 \times 4^3 \\ &= (4 \times 4 \times 4) \times (4 \times 4 \times 4); \\ &= 4^6; \end{aligned}$$

and in the same manner,

$$(a^b)^n = a^{bn};$$

that is, a power a^b is raised to the n^{th} power, by multiplying its exponent. Thus also,

$$(a^b)^n = (a^n)^b;$$

that is, the order of involution is immaterial.

Therefore, multiple involution leads to no further algebraic operations.

8. $4^3 = 64;$

that is, the product of 3 equal factors 4, gives 64.

Inversely, the problem may be, to resolve 64 into a product of 3 equal factors. Each of the factors then will be 4. This reverse operation of involution is called evolution, and is written thus,

$$\sqrt[3]{64} = 4;$$

or, more general,

$$\sqrt[b]{c} = a.$$

$\sqrt[b]{c}$ thus is defined as that number a , which, raised to the power b , gives c ; or, in other words,

$$(\sqrt[b]{c})^b = c.$$

Involution thus far was defined only for integer positive and negative exponents, and the question arises, whether powers with fractional exponents, as $c^{\frac{1}{b}}$ or $c^{\frac{n}{b}}$, have any meaning. Writing,

$$(c^{\frac{1}{b}})^b = c^{b \times \frac{1}{b}} = c^1 = c,$$

it is seen that $c^{\frac{1}{b}}$ is that number, which raised to the power b , gives c ; that is, $c^{\frac{1}{b}}$ is $\sqrt[b]{c}$, and the operation of evolution thus can be expressed as involution with fractional exponent,

$$c^{\frac{1}{b}} = \sqrt[b]{c},$$

and

$$c^{\frac{n}{b}} = (c^{\frac{1}{b}})^n = (\sqrt[b]{c})^n,$$

or,

$$c^{\frac{n}{b}} = (c^n)^{\frac{1}{b}} = \sqrt[b]{c^n},$$

and

$$(\sqrt[b]{c})^n = \sqrt[b]{c^n}.$$

Obviously then,

$$\sqrt[b]{c} = c^{\frac{1}{b}}, \quad \sqrt[b]{c} = c^{\frac{n}{b}},$$

$$\sqrt[b]{c} = c^{-\frac{1}{b}} = \frac{1}{c^{\frac{1}{b}}} = \frac{1}{\sqrt[b]{c}}.$$

Irrational Numbers.

9. Involution with integer exponents, as $4^3 = 64$, can always be carried out. In many cases, evolution can also be carried out. For instance,

$$\sqrt[3]{64} = 4,$$

$$\sqrt[4]{1} = 1;$$

while, in other cases, it cannot be carried out. For instance,

$$\sqrt[3]{2} = ?.$$

Attempting to calculate $\sqrt[4]{2}$, we get,

$$\sqrt[4]{2} = 1.4142135 \dots,$$

and find, no matter how far we carry the calculation, we never come to an end, but get an endless decimal fraction; that is, no number exists in our system of numbers, which can express $\sqrt[4]{2}$, but we can only approximate it, and carry the approximation to any desired degree; some such numbers, as π , have been calculated up to several hundred decimals.

Such numbers as $\sqrt[4]{2}$, which cannot be expressed in any finite form, but merely approximated, are called *irrational numbers*. The name is just as wrong as the name negative number, or imaginary number. There is nothing irrational about $\sqrt[4]{2}$. If we draw a square, with 1 foot as side, the length of the diagonal is $\sqrt[4]{2}$ feet, and the length of the diagonal of a square obviously is just as rational as the length of the sides.

Irrational numbers thus are those real and existing numbers, which cannot be expressed by an integer, or a fraction or finite decimal fraction, but give an endless decimal fraction, which does not repeat.

Endless decimal fractions frequently are met when expressing common fractions as decimals. These decimal representations of common fractions, however, are *periodic* decimals, that is, the numerical values periodically repeat, and in this respect are different from the irrational number, and can, due to their periodic nature, be converted into a finite common fraction. For instance, $2.1387387 \dots$.

Let

$$x = 2.1387387 \dots;$$

then,

$$1000x = 2138.7387387 \dots,$$

subtracting,

$$999x = 2136.6$$

Hence,

$$x = \frac{2136.6}{999} = \frac{21366}{9990} = \frac{1187}{555} = 2\frac{77}{555}.$$

Quadrature Numbers.

10. It is

$$\sqrt[3]{+4} = (+2),$$

since,

$$(+2) \times (+2) = (+4);$$

but it also is:

$$\sqrt[3]{+4} = (-2),$$

since,

$$(-2) \times (-2) = (+4).$$

Therefore, $\sqrt[3]{+4}$ has two values, $(+2)$ and (-2) , and in evolution we thus first strike the interesting feature, that one and the same operation, with the same numerical values, gives several different results.

Since all the positive and negative numbers are used up as the square roots of positive numbers, the question arises, What is the square root of a negative number? For instance, $\sqrt{-4}$ cannot be -2 , as -2 squared gives $+4$, nor can it be $+2$.

$\sqrt{-4} = \sqrt{4 \times (-1)} = \pm 2\sqrt{-1}$, and the question thus resolves itself into: What is $\sqrt{-1}$?

We have derived the absolute numbers from experience, for instance, by measuring distances on a line Fig. 4, from a starting point A .



FIG. 4. Negative and Positive Numbers.

Then we have seen that we get the same distance from A , twice, once toward the right, once toward the left, and this has led to the subdivision of the numbers into positive and negative numbers. Choosing the positive toward the right, in Fig. 4, the negative number would be toward the left (or inversely, choosing the positive toward the left, would give the negative toward the right).

If then we take a number, as $+2$, which represents a distance AB , and multiply by (-1) , we get the distance $AC = -2$

in opposite direction from A . Inversely, if we take $AC = -2$, and multiply by (-1) , we get $AB = +2$; that is, multiplication by (-1) reverses the direction, turns it through 180 deg.

If we multiply $+2$ by $\sqrt{-1}$, we get $+2\sqrt{-1}$, a quantity of which we do not yet know the meaning. Multiplying once more by $\sqrt{-1}$, we get $+2 \times \sqrt{-1} \times \sqrt{-1} = -2$; that is, multiplying a number $+2$, twice by $\sqrt{-1}$, gives a rotation of 180 deg., and multiplication by $\sqrt{-1}$ thus means rotation by half of 180 deg.; or, by 90 deg., and $+2\sqrt{-1}$ thus is the dis-

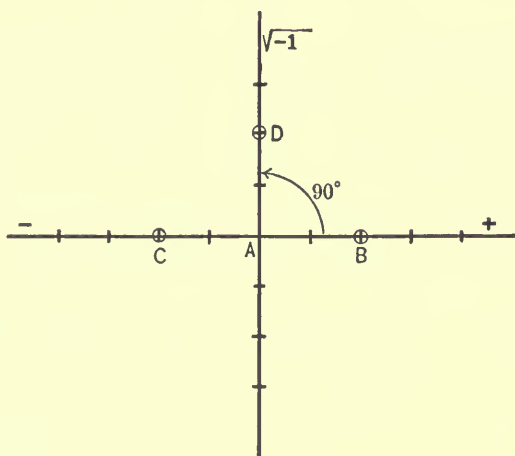


FIG. 5.

tance in the direction rotated 90 deg. from $+2$, or in quadrature direction AD in Fig. 5, and such numbers as $+2\sqrt{-1}$ thus are *quadrature numbers*, that is, represent direction not toward the right, as the positive, nor toward the left, as the negative numbers, but upward or downward.

For convenience of writing, $\sqrt{-1}$ is usually denoted by the letter j .

II. Just as the operation of subtraction introduced in the negative numbers a new kind of numbers, having a direction 180 deg. different, that is, in opposition to the positive numbers, so the operation of evolution introduces in the quadrature number, as $2j$, a new kind of number, having a direction 90 deg.

different; that is, at right angles to the positive and the negative numbers, as illustrated in Fig. 6.

As seen, mathematically the quadrature number is just as real as the negative, physically sometimes the negative number has a meaning—if two opposite directions exist—; sometimes it has no meaning—where one direction only exists. Thus also the quadrature number sometimes has a physical meaning, in those cases where four directions exist, and has no meaning, in those physical problems where only two directions exist.

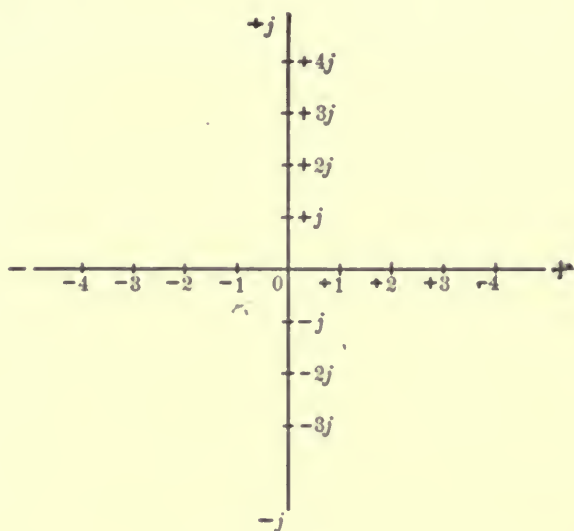


FIG. 6.

For instance, in problems dealing with plain geometry, as in electrical engineering when discussing alternating current vectors in the plane, the quadrature numbers represent the vertical, the ordinary numbers the horizontal direction, and then the one horizontal direction is positive, the other negative, and in the same manner the one vertical direction is positive, the other negative. Usually positive is chosen to the right and upward, negative to the left and downward, as indicated in Fig. 6. In other problems, as when dealing with time, which has only two directions, past and future, the quadrature numbers are not applicable, but only the positive and negative

numbers. In still other problems, as when dealing with illumination, or with individuals, the negative numbers are not applicable, but only the absolute or positive numbers.

Just as multiplication by the negative unit (-1) means rotation by 180° , or reverse of direction, so multiplication by the quadrature unit, j , means rotation by 90° , or change from the horizontal to the vertical direction, and inversely.

General Numbers.

12. By the positive and negative numbers, all the points of a line could be represented numerically as distances from a chosen point A .

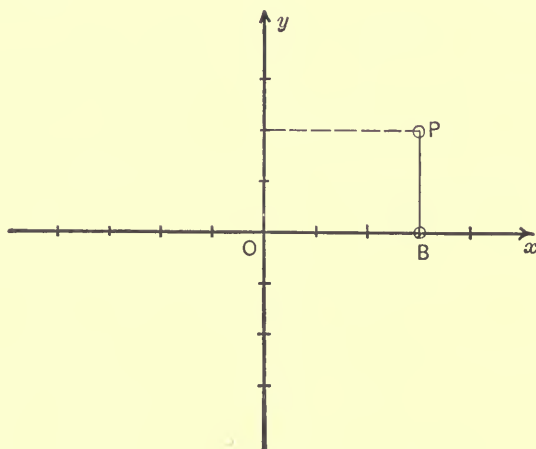


FIG. 7. Simple Vector Diagram.

By the addition of the quadrature numbers, all points of the entire plane can now be represented as distances from chosen coordinate axes x and y , that is, any point P of the plane, Fig. 7, has a horizontal distance, $\overline{OB} = +3$, and a vertical distance, $\overline{BP} = +2j$, and therefore is given by a combination of the distances, $\overline{OB} = +3$ and $\overline{BP} = +2j$. For convenience, the act of combining two such distances in quadrature with each other can be expressed by the plus sign, and the result of combination thereby expressed by $\overline{OB} + \overline{BP} = 3 + 2j$.

Such a combination of an ordinary number and a quadrature number is called a *general number* or a *complex quantity*.

The quadrature number $j b$ thus enormously extends the field of usefulness of algebra, by affording a numerical representation of two-dimensional systems, as the plane, by the general number $a + j b$. They are especially useful and important in electrical engineering, as most problems of alternating currents lead to vector representations in the plane, and therefore can be represented by the general number $a + j b$; that is, the combination of the ordinary number or horizontal distance a , and the quadrature number or vertical distance $j b$.

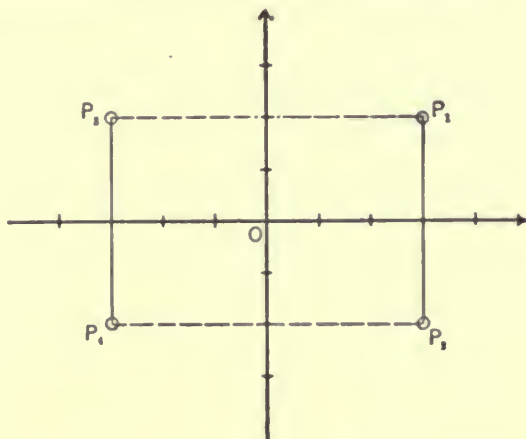


FIG. 8. Vector Diagram.

Analytically, points in the plane are represented by their two coordinates: the horizontal coordinate, or abscissa x , and the vertical coordinate, or ordinate y . Algebraically, in the general number $a + j b$ both coordinates are combined, a being the x coordinate, $j b$ the y coordinate.

Thus in Fig. 8, coordinates of the points are,

$$P_1: x = +3, \quad y = +2 \qquad P_2: x = +3 \quad y = -2,$$

$$P_3: x = -3, \quad y = +2 \qquad P_4: x = -3 \quad y = -2,$$

and the points are located in the plane by the numbers:

$$P_1 = 3 + 2j \quad P_2 = 3 - 2j \quad P_3 = -3 + 2j \quad P_4 = -3 - 2j$$

13. Since already the square root of negative numbers has extended the system of numbers by giving the quadrature number, the question arises whether still further extensions of the system of numbers would result from higher roots of negative quantities.

For instance,

$$\sqrt[4]{-1}=?$$

The meaning of $\sqrt[4]{-1}$ we find in the same manner as that of $\sqrt{-1}$.

A positive number a may be represented on the horizontal axis as P .

Multiplying a by $\sqrt[4]{-1}$ gives $a\sqrt[4]{-1}$, whose meaning we do not yet know. Multiplying again and again by $\sqrt[4]{-1}$, we get, after four multiplications, $a(\sqrt[4]{-1})^4 = -a$; that is, in four steps we have been carried from a to $-a$, a rotation of 180 deg., and

$\sqrt[4]{-1}$ thus means a rotation of $\frac{180}{4} = 45$ deg., therefore, $a\sqrt[4]{-1}$ is the point P_1 in Fig. 9, at distance a from the coordinate center, and under angle 45 deg., which has the coordinates, $x = \frac{a}{\sqrt{2}}$ and $y = \frac{a}{\sqrt{2}}j$; or, is represented by the general number,

$$P_1 = a \frac{1+j}{\sqrt{2}}.$$

$\sqrt[4]{-1}$, however, may also mean a rotation by 135 deg. to P_2 , since this, repeated four times, gives $4 \times 135 = 540$ deg., or the same as 180 deg., or it may mean a rotation by 225 deg. or by 315 deg. Thus four points exist, which represent $a\sqrt[4]{-1}$; the points:

$$P_1 = \frac{+1+j}{\sqrt{2}} a, \quad P_2 = \frac{-1+j}{\sqrt{2}} a,$$

$$P_3 = \frac{-1-j}{\sqrt{2}} a, \quad P_4 = \frac{+1-j}{\sqrt{2}} a.$$

Therefore, $\sqrt[4]{-1}$ is still a general number, consisting of an ordinary and a quadrature number, and thus does not extend our system of numbers any further.

In the same manner, $\sqrt[n]{+1}$ can be found; it is that number, which, multiplied n times with itself, gives $+1$. Thus it represents a rotation by $\frac{360}{n}$ deg., or any multiple thereof; that is, the x coordinate is $\cos q \times \frac{360}{n}$, the y coordinate $\sin q \times \frac{360}{n}$, and,

$$\sqrt[n]{+1} = \cos q \times \frac{360}{n} + j \sin q \times \frac{360}{n},$$

where q is any integer number.

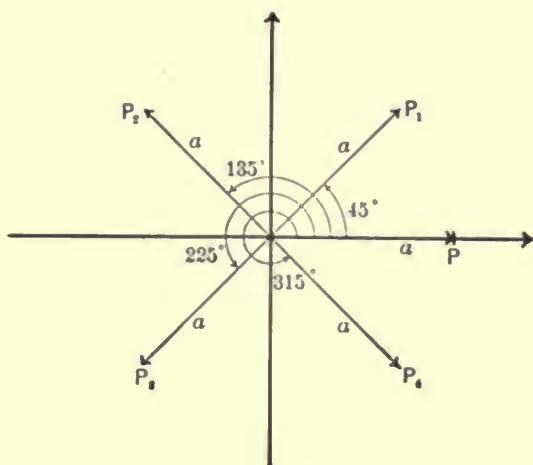


FIG. 9. Vector Diagram $a\sqrt[n]{-1}$.

There are therefore n different values of $a\sqrt[n]{+1}$, which lie equidistant on a circle with radius 1, as shown for $n=9$ in Fig. 10.

14. In the operation of addition, $a+b=c$, the problem is, a and b being given, to find c .

The terms of addition, a and b , are interchangeable, or equivalent, thus: $a+b=b+a$, and addition therefore has only one reverse operation, subtraction; c and b being given, a is found, thus: $a=c-b$, and c and a being given, b is found, thus: $b=c-a$. Either leads to the same operation—subtraction.

The same is the case in multiplication; $a \times b = c$. The

factors a and b are interchangeable or equivalent; $a \times b = b \times a$ and the reverse operation, division, $a = \frac{c}{b}$ is the same as $b = \frac{c}{a}$.

In involution, however, $a^b = c$, the two numbers a and b are not interchangeable, and a^b is not equal to b^a . For instance $4^3 = 64$ and $3^4 = 81$.

Therefore, involution has two reverse operations:

(a) c and b given, a to be found,

$$a = \sqrt[b]{c};$$

or evolution,

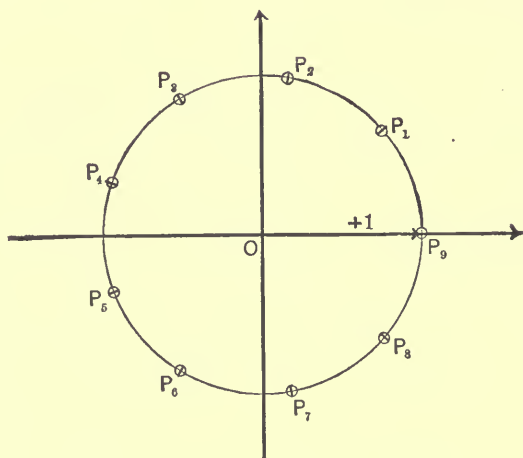


FIG. 10. Points Determined by $\sqrt[n]{+1}$.

(b) c and a given, b to be found,

$$b = \log_a c;$$

or, logarithmation.

Logarithmation.

15. Logarithmation thus is one of the reverse operations of involution, and the logarithm is the exponent of involution.

Thus a logarithmic expression may be changed to an exponential, and inversely, and the laws of logarithmation are the laws, which the exponents obey in involution.

1. Powers of equal base are multiplied by adding the exponents: $a^b \times a^n = a^{b+n}$. Therefore, the logarithm of a

product is the sum of the logarithms of the factors, thus $\log_a c \times d = \log_a c + \log_a d$.

2. A power is raised to a power by multiplying the exponents: $(a^b)^n = a^{bn}$.

Therefore the logarithm of a power is the exponent times the logarithm of the base, or, the number under the logarithm is raised to the power n , by multiplying the logarithm by n :

$$\log_a c^n = n \log_a c,$$

$\log_a 1 = 0$, because $a^0 = 1$. If the base $a > 1$, $\log_a c$ is positive, if $c > 1$, and is negative, if $c < 1$, but > 0 . The reverse is the case, if $a < 1$. Thus, the logarithm traverses all positive and negative values for the positive values of c , and the logarithm of a negative number thus can be neither positive nor negative.

$\log_a (-c) = \log_a c + \log_a (-1)$, and the question of finding the logarithms of negative numbers thus resolves itself into finding the value of $\log_a (-1)$.

There are two standard systems of logarithms one with the base $e = 2.71828 \dots^*$, and the other with the base 10 is used, the former in algebraic, the latter in numerical calculations. Logarithms of any base a can easily be reduced to any other base.

For instance, to reduce $b = \log_a c$ to the base 10: $b = \log_a c$ means, in the form of involution: $a^b = c$. Taking the logarithm hereof gives, $b \log_{10} a = \log_{10} c$, hence,

$$b = \frac{\log_{10} c}{\log_{10} a}; \quad \text{or} \quad \log_a c = \frac{\log_{10} c}{\log_{10} a}.$$

Thus, regarding the logarithms of negative numbers, we need to consider only $\log_{10} (-1)$ or $\log_e (-1)$.

If $jx = \log_e (-1)$, then $e^{jx} = -1$,

and since, as will be seen in Chapter II,

$$e^{jx} = \cos x + j \sin x,$$

it follows that,

$$\cos x + j \sin x = -1.$$

* Regarding e , see Chapter II, p. 71.

Hence, $x = \pi$, or an odd multiple thereof, and

$$\log_e(-1) = j\pi(2n+1),$$

where n is any integer number.

Thus logarithmation also leads to the quadrature number j , but to no further extension of the system of numbers.

Quaternions.

16. Addition and subtraction, multiplication and division, involution and evolution and logarithmation thus represent all the algebraic operations, and the system of numbers in which all these operations can be carried out under all conditions is that of the general number, $a + jb$, comprising the ordinary number a and the quadrature number jb . The number a as well as b may be positive or negative, may be integer, fraction or irrational.

Since by the introduction of the quadrature number jb , the application of the system of numbers was extended from the line, or more general, one-dimensional quantity, to the plane, or the two-dimensional quantity, the question arises, whether the system of numbers could be still further extended, into three dimensions, so as to represent space geometry. While in electrical engineering most problems lead only to plain figures, vector diagrams in the plane, occasionally space figures would be advantageous if they could be expressed algebraically. Especially in mechanics this would be of importance when dealing with forces as vectors in space.

In the quaternion calculus methods have been devised to deal with space problems. The quaternion calculus, however, has not yet found an engineering application comparable with that of the general number, or, as it is frequently called, the *complex quantity*. The reason is that the quaternion is not an algebraic quantity, and the laws of algebra do not uniformly apply to it.

17. With the rectangular coordinate system in the plane, Fig. 11, the x axis may represent the ordinary numbers, the y axis the quadrature numbers, and multiplication by $j = \sqrt{-1}$ represents rotation by 90 deg. For instance, if P_1 is a point

$a+jb=3+2j$, the point P_2 , 90 deg. away from P_1 , would be:

$$P_2 = jP_1 = j(a+jb) = j(3+2j) = -2+3j,$$

To extend into space, we have to add the third or z axis, as shown in perspective in Fig. 12. Rotation in the plane xy , by 90 deg., in the direction $+x$ to $+y$, then means multiplication by j . In the same manner, rotation in the yz plane, by 90 deg., from $+y$ to $+z$, would be represented by multiplica-

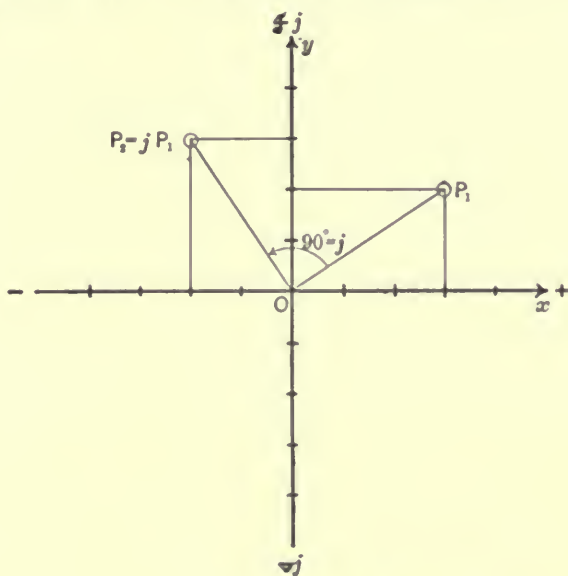


FIG. 11. Vectors in a Plane.

tion with h , and rotation by 90 deg. in the zx plane, from $+z$ to $+x$ would be presented by k , as indicated in Fig. 12.

All three of these rotors, j , h , k , would be $\sqrt{-1}$, since each, applied twice, reverses the direction, that is, represents multiplication by (-1) .

As seen in Fig. 12, starting from $+x$, and going to $+y$, then to $+z$, and then to $+x$, means successive multiplication by j , h and k , and since we come back to the starting point, the total operation produces no change, that is, represents multiplication by $(+1)$. Hence, it must be,

$$j h k = +1.$$

Algebraically this is not possible, since each of the three quantities is $\sqrt{-1}$, and $\sqrt{-1} \times \sqrt{-1} \times \sqrt{-1} = -\sqrt{-1}$, and not $(+1)$.

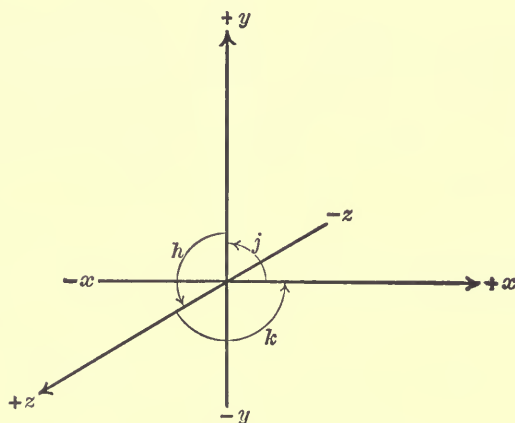


FIG. 12. Vectors in Space, $jhk = +1$.

If we now proceed again from x , in positive rotation, but first turn in the xz plane, we reach by multiplication with k the negative z axis, $-z$, as seen in Fig. 13. Further multiplica-

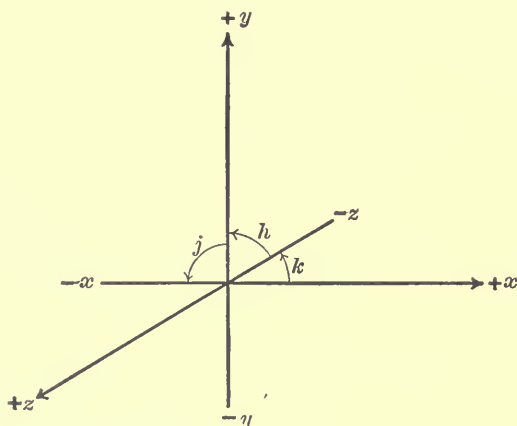


FIG. 13. Vectors in Space, $khj = -1$.

tion by h brings us to $+y$, and multiplication by j to $-x$, and in this case the result of the three successive rotations by

90 deg., in the same direction as in Fig. 12, but in a different order, is a reverse; that is, represents (-1) . Therefore,

$$khj = -1,$$

and hence,

$$jkk = -khj.$$

Thus, in vector analysis of space, we see that the fundamental law of algebra,

$$a \times b = b \times a,$$

does not apply, and the order of the factors of a product is not immaterial, but by changing the order of the factors of the product jkk , its sign was reversed. Thus common factors cannot be canceled as in algebra; for instance, if in the correct expression, $jkk = -khj$, we should cancel by j , h and k , as could be done in algebra, we would get $+1 = -1$, which is obviously wrong.

For this reason all the mechanisms devised for vector analysis in space have proven more difficult in their application, and have not yet been used to any great extent in engineering practice.

B. ALGEBRA OF THE GENERAL NUMBER, OR COMPLEX QUANTITY.

Rectangular and Polar Coordinates.

18. The general number, or complex quantity, $a + jb$, is the most general expression to which the laws of algebra apply. It therefore can be handled in the same manner and under the same rules as the ordinary number of elementary arithmetic. The only feature which must be kept in mind is that $j^2 = -1$, and where in multiplication or other operations j^2 occurs, it is replaced by its value, -1 . Thus, for instance,

$$\begin{aligned}(a + jb)(c + jd) &= ac + jad + jbc + j^2bd \\ &= ac + jad + jbc - bd \\ &= (ac - bd) + j(ad + bc).\end{aligned}$$

Herefrom it follows that all the higher powers of j can be eliminated, thus:

$$\begin{aligned}j^1 &= j, & j^2 &= -1, & j^3 &= -j, & j^4 &= +1; \\ j^5 &= +j, & j^6 &= -1, & j^7 &= -j, & j^8 &= +1; \\ j^9 &= +j, & & & & & \text{etc.}\end{aligned}$$

In distinction from the general number or complex quantity, the ordinary numbers, $+a$ and $-a$, are occasionally called *scalars*, or *real* numbers. The general number thus consists of the combination of a scalar or real number and a quadrature number, or imaginary number.

Since a quadrature number cannot be equal to an ordinary number it follows that, if two general numbers are equal, their real components or ordinary numbers, as well as their quadrature numbers or imaginary components must be equal, thus, if

$$a + jb = c + jd,$$

then,

$$a = c \quad \text{and} \quad b = d.$$

Every equation with general numbers thus can be resolved into two equations, one containing only the ordinary numbers, the other only the quadrature numbers. For instance, if

$$x + jy = 5 - 3j,$$

then,

$$x = 5 \quad \text{and} \quad y = -3.$$

19. The best way of getting a conception of the general number, and the algebraic operations with it, is to consider the general number as representing a point in the plane. Thus the general number $a + jb = 6 + 2.5j$ may be considered as representing a point P , in Fig. 14, which has the horizontal distance from the y axis, $\overline{OA} = \overline{BP} = a = 6$, and the vertical distance from the x axis, $\overline{OB} = \overline{AP} = b = 2.5$.

The total distance of the point P from the coordinate center O then is

$$\begin{aligned} \overline{OP} &= \sqrt{\overline{OA}^2 + \overline{AP}^2} \\ &= \sqrt{a^2 + b^2} = \sqrt{6^2 + 2.5^2} = 6.5, \end{aligned}$$

and the angle, which this distance OP makes with the x axis, is given by

$$\begin{aligned} \tan \theta &= \frac{\overline{AP}}{\overline{OA}} \\ &= \frac{b}{a} = \frac{2.5}{6} = 0.417. \end{aligned}$$

Instead of representing the general number by the two components, a and b , in the form $a+jb$, it can also be represented by the two quantities: the distance of the point P from the center O ,

$$c = \sqrt{a^2 + b^2};$$

and the angle between this distance and the x axis,

$$\tan \theta = \frac{b}{a}.$$

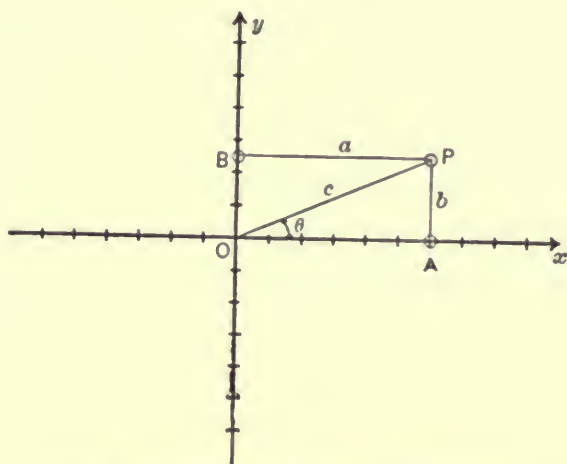


FIG. 14. Rectangular and Polar Coordinates.

Then referring to Fig. 14,

$$a = c \cos \theta \quad \text{and} \quad b = c \sin \theta,$$

and the general number $a+jb$ thus can also be written in the form,

$$c(\cos \theta + j \sin \theta).$$

The form $a+jb$ expresses the general number by its rectangular components a and b , and corresponds to the rectangular coordinates of analytic geometry; a is the x coordinate, b the y coordinate.

The form $c(\cos \theta + j \sin \theta)$ expresses the general number by what may be called its polar components, the radius c and the

angle θ , and corresponds to the polar coordinates of analytic geometry. c is frequently called the radius vector or scalar, θ the phase angle of the general number.

While usually the rectangular form $a+jb$ is more convenient, sometimes the polar form $c(\cos \theta + j \sin \theta)$ is preferable, and transformation from one form to the other therefore frequently applied.

Addition and Subtraction.

20. If $a_1 + jb_1 = 6 + 2.5j$ is represented by the point P_1 ; this point is reached by going the horizontal distance $a_1 = 6$ and the vertical distance $b_1 = 2.5$. If $a_2 + jb_2 = 3 + 4j$ is represented by the point P_2 , this point is reached by going the horizontal distance $a_2 = 3$ and the vertical distance $b_2 = 4$.

The sum of the two general numbers $(a_1 + jb_1) + (a_2 + jb_2) = (6 + 2.5j) + (3 + 4j)$, then is given by point P_0 , which is reached by going a horizontal distance equal to the sum of the horizontal distances of P_1 and P_2 : $a_0 = a_1 + a_2 = 6 + 3 = 9$, and a vertical distance equal to the sum of the vertical distances of P_1 and P_2 : $b_0 = b_1 + b_2 = 2.5 + 4 = 6.5$, hence, is given by the general number

$$\begin{aligned} a_0 + jb_0 &= (a_1 + a_2) + j(b_1 + b_2) \\ &= 9 + 6.5j. \end{aligned}$$

Geometrically, point P_0 is derived from points P_1 and P_2 , by the diagonal $\overline{OP_0}$ of the parallelogram $OP_1P_0P_2$, constructed with $\overline{OP_1}$ and $\overline{OP_2}$ as sides, as seen in Fig. 15.

Herefrom it follows that addition of general numbers represents geometrical combination by the parallelogram law.

Inversely, if P_0 represents the number

$$a_0 + jb_0 = 9 + 6.5j,$$

and P_1 represents the number

$$a_1 + jb_1 = 6 + 2.5j,$$

the difference of these numbers will be represented by a point P_2 , which is reached by going the difference of the horizontal

distances and of the vertical distances of the points P_0 and P_1 . P_2 thus is represented by

$$a_2 = a_0 - a_1 = 9 - 6 = 3,$$

and

$$b_2 = b_0 - b_1 = 6.5 - 2.5 = 4.$$

Therefore, the difference of the two general numbers $(a_0 + jb_0)$ and $(a_1 + jb_1)$ is given by the general number:

$$\begin{aligned} a_2 + jb_2 &= (a_0 - a_1) + j(b_0 - b_1) \\ &= 3 + 4j, \end{aligned}$$

as seen in Fig. 15.

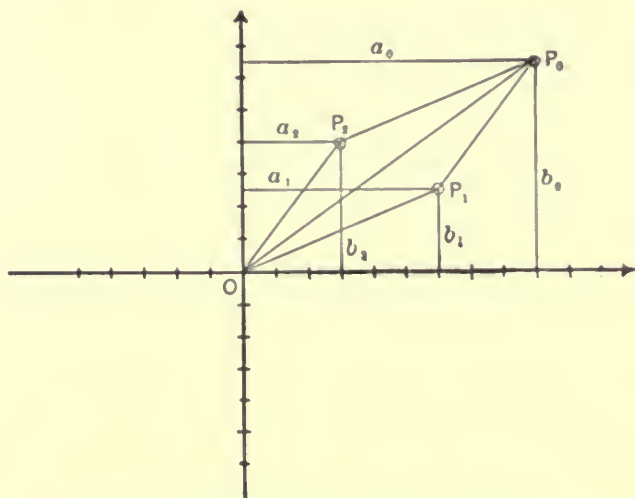


FIG. 15. Addition and Subtraction of Vectors.

This difference $a_2 + jb_2$ is represented by one side $\overline{OP_2}$ of the parallelogram $OP_1P_0P_2$, which has $\overline{OP_1}$ as the other side, and $\overline{OP_0}$ as the diagonal.

Subtraction of general numbers thus geometrically represents the resolution of a vector $\overline{OP_0}$ into two components $\overline{OP_1}$ and $\overline{OP_2}$, by the parallelogram law.

Herein lies the main advantage of the use of the general number in engineering calculation: If the vectors are represented by general numbers (complex quantities), combination and resolution of vectors by the parallelogram law is carried out by

simple addition or subtraction of their general numerical values, that is, by the simplest operation of algebra.

21. General numbers are usually denoted by capitals, and their rectangular components, the ordinary number and the quadrature number, by small letters, thus:

$$A = a_1 + ja_2;$$

the distance of the point which represents the general number A from the coordinate center is called the *absolute value*, radius or scalar of the general number or complex quantity. It is the vector a in the polar representation of the general number:

$$A = a(\cos \theta + j \sin \theta),$$

and is given by $a = \sqrt{a_1^2 + a_2^2}$.

The absolute value, or scalar, of the general number is usually also denoted by small letters, but sometimes by capitals, and in the latter case it is distinguished from the general number by using a different type for the latter, or underlining or dotting it, thus:

$$\begin{array}{lll} A = a_1 + ja_2; & \text{or} & \dot{A} = a_1 + ja_2; \text{ or } \underline{A} = a_1 + ja_2 \\ \text{or} & \dot{A} = a_1 + ja_2; & \text{or} & A = a_1 + ja_2 \\ & a = \sqrt{a_1^2 + a_2^2}; & \text{or} & A = \sqrt{a_1^2 + a_2^2}, \end{array}$$

$$\text{and} \quad a_1 + ja_2 = a(\cos \theta + j \sin \theta);$$

$$\text{or} \quad a_1 + ja_2 = A(\cos \theta + j \sin \theta).$$

22. The absolute value, or scalar, of a general number is always an absolute number, and positive, that is, the sign of the rectangular component is represented in the angle θ . Thus referring to Fig. 16,

$$\dot{A} = a_1 + ja_2 = 4 + 3j;$$

$$\text{gives,} \quad a = \sqrt{a_1^2 + a_2^2} = 5;$$

$$\tan \theta = \frac{3}{4} = 0.75;$$

$$\theta = 37 \text{ deg.};$$

$$\text{and} \quad \dot{A} = 5 (\cos 37 \text{ deg.} + j \sin 37 \text{ deg}).$$

The expression

$$A = a_1 + ja_2 = 4 - 3j$$

gives

$$a = \sqrt{a_1^2 + a_2^2} = 5;$$

$$\tan \theta = -\frac{3}{4} = -0.75;$$

$$\theta = -37 \text{ deg.}; \text{ or } = 180 - 37 = 143 \text{ deg.}$$

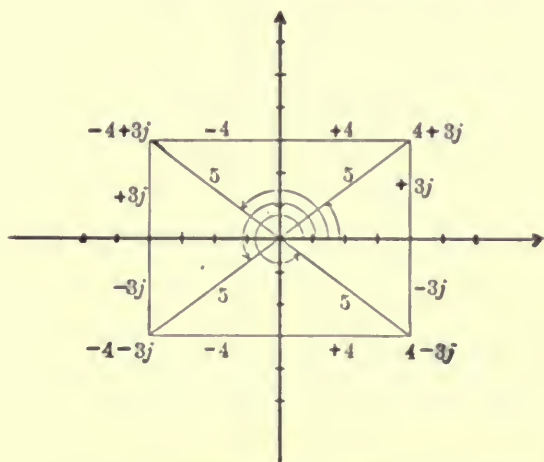


FIG. 16. Representation of General Numbers.

Which of the two values of θ is the correct one is seen from the condition $a_1 = a \cos \theta$. As a_1 is positive, $+4$, it follows that $\cos \theta$ must be positive; $\cos (-37 \text{ deg.})$ is positive, $\cos 143 \text{ deg.}$ is negative; hence the former value is correct:

$$\begin{aligned} A &= 5\{\cos(-37 \text{ deg.}) + j \sin(-37 \text{ deg.})\} \\ &= 5(\cos 37 \text{ deg.} - j \sin 37 \text{ deg.}). \end{aligned}$$

Two such general numbers as $(4 + 3j)$ and $(4 - 3j)$, or, in general,

$$(a + jb) \text{ and } (a - jb),$$

are called conjugate numbers. Their product is an ordinary and not a general number, thus: $(a + jb)(a - jb) = a^2 + b^2$.

The expression

$$\dot{A} = a_1 + ja_2 = -4 + 3j$$

gives

$$a = \sqrt{a_1^2 + a_2^2} = 5;$$

$$\tan \theta = -\frac{3}{4} = -0.75;$$

$$\theta = -37 \text{ deg.} \quad \text{or} \quad = 180 - 37 = 143 \text{ deg.};$$

but since $a_1 = a \cos \theta$ is negative, -4 , $\cos \theta$ must be negative, hence, $\theta = 143 \text{ deg.}$ is the correct value, and

$$\begin{aligned} \dot{A} &= 5(\cos 143 \text{ deg.} + j \sin 143 \text{ deg.}) \\ &= 5(-\cos 37 \text{ deg.} + j \sin 37 \text{ deg.}) \end{aligned}$$

The expression

$$\dot{A} = a_1 + ja_2 = -4 - 3j,$$

gives

$$a = \sqrt{a_1^2 + a_2^2} = 5;$$

$$\theta = 37 \text{ deg.}; \quad \text{or} \quad = 180 + 37 = 217 \text{ deg.};$$

but since $a_1 = a \cos \theta$ is negative, -4 , $\cos \theta$ must be negative, hence $\theta = 217 \text{ deg.}$ is the correct value, and,

$$\begin{aligned} \dot{A} &= 5 (\cos 217 \text{ deg.} + j \sin 217 \text{ deg.}) \\ &= 5(-\cos 37 \text{ deg.} - j \sin 37 \text{ deg.}) \end{aligned}$$

The four general numbers, $+4 + 3j$, $+4 - 3j$, $-4 + 3j$, and $-4 - 3j$, have the same absolute value, 5, and in their representations as points in a plane have symmetrical locations in the four quadrants, as shown in Fig. 16.

As the general number $\dot{A} = a_1 + ja_2$ finds its main use in representing vectors in the plane, it very frequently is called a vector quantity, and the algebra of the general number is spoken of as *vector analysis*.

Since the general numbers $\dot{A} = a_1 + ja_2$ can be made to represent the points of a plane, they also may be called *plane numbers*, while the positive and negative numbers, $+a$ and $-a$,

may be called the *linear numbers*, as they represent the points of a line.

Example: Steam Path in a Turbine.

23. As an example of a simple operation with general numbers one may calculate the steam path in a two-wheel stage of an impulse turbine.

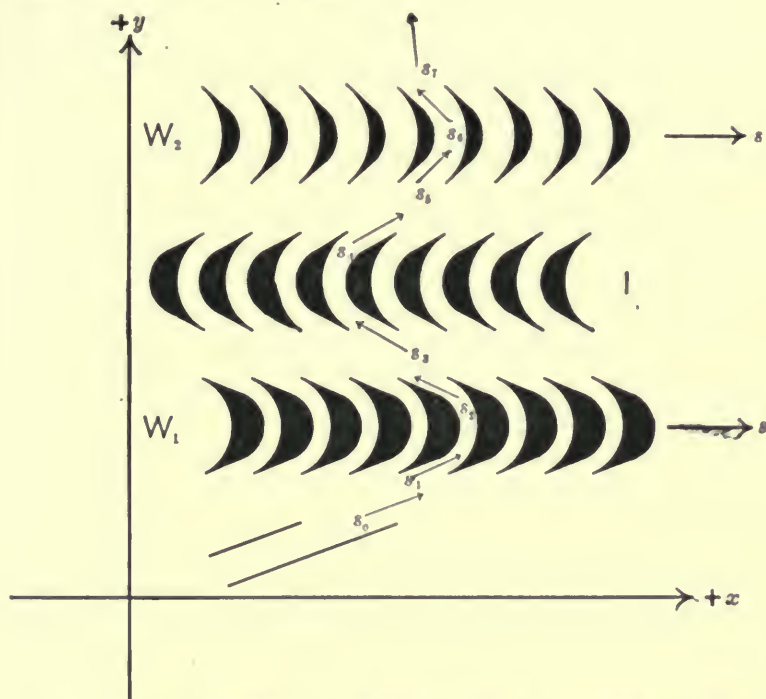


FIG. 17. Path of Steam in a Two-wheel Stage of an Impulse Turbine.

Let Fig. 17 represent diagrammatically a tangential section through the bucket rings of the turbine wheels. W_1 and W_2 are the two revolving wheels, moving in the direction indicated by the arrows, with the velocity $s = 400$ feet per sec. I are the stationary intermediate buckets, which turn the exhaust steam from the first bucket wheel W_1 , back into the direction required to impinge on the second bucket wheel W_2 . The steam jet issues from the expansion nozzle at the speed $s_0 = 2200$

feet per sec., and under the angle $\theta_0 = 20$ deg., against the first bucket wheel W_1 .

The exhaust angles of the three successive rows of buckets, W_1 , I , and W_2 , are respectively 24 deg., 30 deg. and 45 deg. These angles are calculated from the section of the bucket exit required to pass the steam at its momentary velocity, and from the height of the passage required to give no steam eddies, in a manner which is of no interest here.

As friction coefficient in the bucket passages may be assumed $k_f = 0.12$; that is, the exit velocity is $1 - k_f = 0.88$ of the entrance velocity of the steam in the buckets.

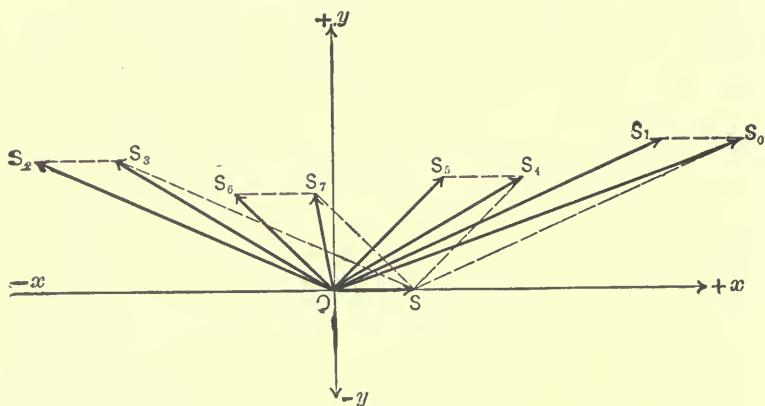


FIG. 18. Vector Diagram of Velocities of Steam in Turbine.

Choosing then as x -axis the direction of the tangential velocity of the turbine wheels, as y -axis the axial direction, the velocity of the steam supply from the expansion nozzle is represented in Fig. 18 by a vector $\overline{OS}_0$ of length $s_0 = 2200$ feet per sec., making an angle $\theta_0 = 20$ deg. with the x -axis; hence, can be expressed by the general number or vector quantity:

$$\begin{aligned}\dot{S}_0 &= s_0 (\cos \theta_0 + j \sin \theta_0) \\ &= 2200 (\cos 20 \text{ deg.} + j \sin 20 \text{ deg.}) \\ &= 2070 + 750j \text{ ft. per sec.}\end{aligned}$$

The velocity of the turbine wheel W_1 is $s = 400$ feet per second, and represented in Fig. 18 by the vector $\overline{OS}$, in horizontal direction.

The relative velocity with which the steam enters the bucket passage of the first turbine wheel W_1 thus is:

$$\begin{aligned}\dot{S}_1 &= \dot{S}_0 - s \\ &= (2070 + 750j) - 400 \\ &= 1670 + 750j \text{ ft. per sec.}\end{aligned}$$

This vector is shown as $\overline{OS}_1$ in Fig. 18.

The angle θ_1 , under which the steam enters the bucket passage thus is given by

$$\tan \theta_1 = \frac{750}{1670} = 0.450, \text{ as } \theta_1 = 24.3 \text{ deg.}$$

This angle thus has to be given to the front edge of the buckets of the turbine wheel W_1 .

The absolute value of the relative velocity of steam jet and turbine wheel W_1 , at the entrance into the bucket passage, is

$$s_1 = \sqrt{1670^2 + 750^2} = 1830 \text{ ft. per sec.}$$

In traversing the bucket passages the steam velocity decreases by friction etc., from the entrance value s_1 to the exit value

$$s_2 = s_1(1 - k_f) = 1830 \times 0.88 = 1610 \text{ ft. per sec.,}$$

and since the exit angle of the bucket passage has been chosen as $\theta_2 = 24$ deg., the relative velocity with which the steam leaves the first bucket wheel W_1 is represented by a vector $\overline{OS}_2$ in Fig. 18, of length $s_2 = 1610$, under angle 24 deg. The steam leaves the first wheel in backward direction, as seen in Fig. 17, and 24 deg. thus is the angle between the steam jet and the negative x -axis; hence, $\theta_2 = 180 - 24 = 156$ deg. is the vector angle. The relative steam velocity at the exit from wheel W_1 can thus be represented by the vector quantity

$$\begin{aligned}\dot{S}_2 &= s_2(\cos \theta_2 + j \sin \theta_2) \\ &= 1610 (\cos 156 \text{ deg.} + j \sin 156 \text{ deg.}) \\ &= -1470 + 655 j.\end{aligned}$$

Since the velocity of the turbine wheel W_1 is $s = 400$, the velocity of the steam in space, after leaving the first turbine

wheel, that is, the velocity with which the steam enters the intermediate I , is

$$\begin{aligned}\dot{S}_3 &= \dot{S}_2 + s \\ &= (-1470 + 655j) + 400 \\ &= -1070 + 655j,\end{aligned}$$

and is represented by vector $\overline{OS}_3$ in Fig. 18. The direction of this steam jet is given by

$$\tan \theta_3 = -\frac{655}{1070} = -0.613,$$

as

$$\theta_3 = -31.6 \text{ deg.}; \text{ or, } 180 - 31.6 = 148.4 \text{ deg.}$$

The latter value is correct, as $\cos \theta_3$ is negative, and $\sin \theta_3$ is positive.

The steam jet thus enters the intermediate under the angle of 148.4 deg.; that is, the angle $180 - 148.4 = 31.6$ deg. in opposite direction. The buckets of the intermediate I thus must be curved in reverse direction to those of the wheel W_1 , and must be given the angle 31.6 deg. at their front edge.

The absolute value of the entrance velocity into the intermediate I is

$$s_3 = \sqrt{1070^2 + 655^2} = 1255 \text{ ft. per sec.}$$

In passing through the bucket passages, this velocity decreases by friction, to the value:

$$s_4 = s_3(1 - k_f) = 1255 \times 0.88 = 1105 \text{ ft. per sec.,}$$

and since the exit edge of the intermediate is given the angle: $\theta_4 = 30$ deg., the exit velocity of the steam from the intermediate is represented by the vector $\overline{OS}_4$ in Fig. 18, of length $s_4 = 1105$, and angle $\theta_4 = 30$ deg., hence,

$$\begin{aligned}\dot{S}_4 &= 1105 (\cos 30 \text{ deg.} + j \sin 30 \text{ deg.}) \\ &= 955 + 550j \text{ ft. per sec.}\end{aligned}$$

This is the velocity with which the steam jet impinges on the second turbine wheel W_2 , and as this wheel revolves

with velocity $s=400$, the relative velocity, that is, the velocity with which the steam enters the bucket passages of wheel W_2 , is,

$$\begin{aligned}\dot{S}_5 &= \dot{S}_4 - s \\ &= (955 + 550j) - 400 \\ &= 555 + 550j \text{ ft. per sec.};\end{aligned}$$

and is represented by vector $\overline{OS}_5$ in Fig. 18.

The direction of this steam jet is given by

$$\tan \theta_5 = \frac{550}{555} = 0.990, \quad \text{as } \theta_5 = 44.8 \text{ deg.}$$

Therefore, the entrance edge of the buckets of the second wheel W_2 must be shaped under angle $\theta_5 = 44.8$ deg.

The absolute value of the entrance velocity is

$$s_5 = \sqrt{555^2 + 550^2} = 780 \text{ ft. per sec.}$$

In traversing the bucket passages, the velocity drops from the entrance value S_5 , to the exit value,

$$s_6 = s_5(1 - k_f) = 780 \times 0.88 = 690 \text{ ft. per sec.}$$

Since the exit angles of the buckets of wheel W_2 has been chosen as 45 deg., and the exit is in backward direction, $\theta_6 = 180 - 45 = 135$ deg., the steam jet velocity at the exit of the bucket passages of the last wheel is given by the general number

$$\begin{aligned}\dot{S}_6 &= s_6(\cos \theta_6 + j \sin \theta_6) \\ &= 690 (\cos 135 \text{ deg.} + j \sin 135 \text{ deg.}) \\ &= -487 + 487j \text{ ft. per sec.,}\end{aligned}$$

and represented by vector $\overline{OS}_6$ in Fig. 18.

Since $s=400$ is the wheel velocity, the velocity of the steam after leaving the last wheel W_2 , that is, the "lost" or "rejected" velocity, is

$$\begin{aligned}\dot{S}_7 &= \dot{S}_6 + s \\ &= (-487 + 487j) + 400 \\ &= -87 + 487j \text{ ft. per sec.,}\end{aligned}$$

and is represented by vector $\overline{OS}_7$ in Fig. 18.

The direction of the exhaust steam is given by,

$$\tan \theta_7 = -\frac{487}{87} = -5.6, \quad \text{as } \theta_7 = 180 - 80 = 100 \text{ deg.},$$

and the absolute velocity is,

$$s_7 = \sqrt{87^2 + 487^2} = 495 \text{ ft. per sec.}$$

Multiplication of General Numbers.

24. If $\dot{A} = a_1 + ja_2$ and $\dot{B} = b_1 + jb_2$, are two general, or plane numbers, their product is given by multiplication, thus:

$$\begin{aligned} \dot{A}\dot{B} &= (a_1 + ja_2)(b_1 + jb_2) \\ &= a_1b_1 + ja_1b_2 + ja_2b_1 + j^2a_2b_2, \end{aligned}$$

and since $j^2 = -1$,

$$\dot{A}\dot{B} = (a_1b_1 - a_2b_2) + j(a_1b_2 + a_2b_1),$$

and the product can also be represented in the plane, by a point,

$$\dot{C} = c_1 + jc_2,$$

where,

$$c_1 = a_1b_1 - a_2b_2,$$

and

$$c_2 = a_1b_2 + a_2b_1.$$

For instance, $\dot{A} = 2 + j$ multiplied by $\dot{B} = 1 + 1.5j$ gives

$$c_1 = 2 \times 1 - 1 \times 1.5 = 0.5,$$

$$c_2 = 2 \times 1.5 + 1 \times 1 = 4;$$

hence,

$$\dot{C} = 0.5 + 4j,$$

as shown in Fig. 19.

25. The geometrical relation between the factors $\dot{A}$ and $\dot{B}$ and the product $\dot{C}$ is better shown by using the polar expression; hence, substituting,

$$\left. \begin{aligned} a_1 &= a \cos \alpha \\ a_2 &= a \sin \alpha \end{aligned} \right\} \quad \text{and} \quad \left. \begin{aligned} b_1 &= b \cos \beta \\ b_2 &= b \sin \beta \end{aligned} \right\},$$

which gives

$$\left. \begin{aligned} a &= \sqrt{a_1^2 + a_2^2} \\ \tan \alpha &= \frac{a_2}{a_1} \end{aligned} \right\} \quad \text{and} \quad \left. \begin{aligned} b &= \sqrt{b_1^2 + b_2^2} \\ \tan \beta &= \frac{b_2}{b_1} \end{aligned} \right\},$$

the quantities may be written thus:

$$\dot{A} = a(\cos \alpha + j \sin \alpha);$$

$$\dot{B} = b(\cos \beta + j \sin \beta),$$

and then,

$$\begin{aligned} \dot{C} = \dot{A}\dot{B} &= ab(\cos \alpha + j \sin \alpha)(\cos \beta + j \sin \beta) \\ &= ab \{(\cos \alpha \cos \beta - \sin \alpha \sin \beta) + j(\cos \alpha \sin \beta + \sin \alpha \cos \beta)\} \\ &= ab \{\cos (\alpha + \beta) + j \sin (\alpha + \beta)\}; \end{aligned}$$

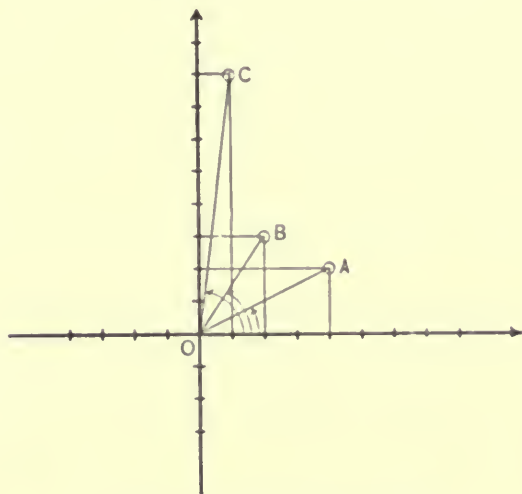


FIG. 19. Multiplication of Vectors.

that is, two general numbers are multiplied by multiplying their absolute values or vectors, a and b , and adding their phase angles α and β .

Thus, to multiply the vector quantity, $\dot{A} = a_1 + ja_2 = a(\cos \alpha + j \sin \alpha)$ by $\dot{B} = b_1 + jb_2 = b(\cos \beta + j \sin \beta)$ the vector OA in Fig. 19, which represents the general number $\dot{A}$, is increased by the factor $b = \sqrt{b_1^2 + b_2^2}$, and rotated by the angle β , which is given by $\tan \beta = \frac{b_2}{b_1}$.

Thus, a complex multiplier $\dot{B}$ turns the direction of the multiplicand $\dot{A}$, by the phase angle of the multiplier $\dot{B}$, and multiplies the absolute value or vector of $\dot{A}$, by the absolute value of $\dot{B}$ as factor.

The multiplier B is occasionally called an *operator*, as it carries out the operation of rotating the direction and changing the length of the multiplicand.

26. In multiplication, division and other algebraic operations with the representations of physical quantities (as alternating currents, voltages, impedances, etc.) by mathematical symbols, whether ordinary numbers or general numbers, it is necessary to consider whether the result of the algebraic operation, for instance, the product of two factors, has a physical meaning, and if it has a physical meaning, whether this meaning is such that the product can be represented in the same diagram as the factors.

For instance, $3 \times 4 = 12$; but 3 horses $\times$ 4 horses does not give 12 horses, nor 12 horses², but is physically meaningless. However, 3 ft. $\times$ 4 ft. = 12 sq.ft. Thus, if the numbers represent

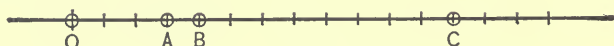


FIG. 20.

horses, multiplication has no physical meaning. If they represent feet, the product of multiplication has a physical meaning, but a meaning which differs from that of the factors. Thus, if on the line in Fig. 20, $\overline{OA} = 3$ feet, $\overline{OB} = 4$ feet, the product, 12 square feet, while it has a physical meaning, cannot be represented any more by a point on the same line; it is not the point $\overline{OC} = 12$, because, if we expressed the distances $\overline{OA}$ and $\overline{OB}$ in inches, 36 and 48 inches respectively, the product would be $36 \times 48 = 1728$ sq.in., while the distance $\overline{OC}$ would be 144 inches.

27. In all mathematical operations with physical quantities it therefore is necessary to consider at every step of the mathematical operation, whether it still has a physical meaning, and, if graphical representation is resorted to, whether the nature of the physical meaning is such as to allow graphical representation in the same diagram, or not.

An instance of this general limitation of the application of mathematics to physical quantities occurs in the representation of alternating current phenomena by general numbers, or complex quantities

An alternating current can be represented by a vector $\overline{OI}$ in a polar diagram, Fig. 21, in which one complete revolution or 360 deg. represents the time of one complete period of the alternating current. This vector $\overline{OI}$ can be represented by a general number,

$$I = i_1 + ji_2,$$

where i_1 is the horizontal, i_2 the vertical component of the current vector $\overline{OI}$.

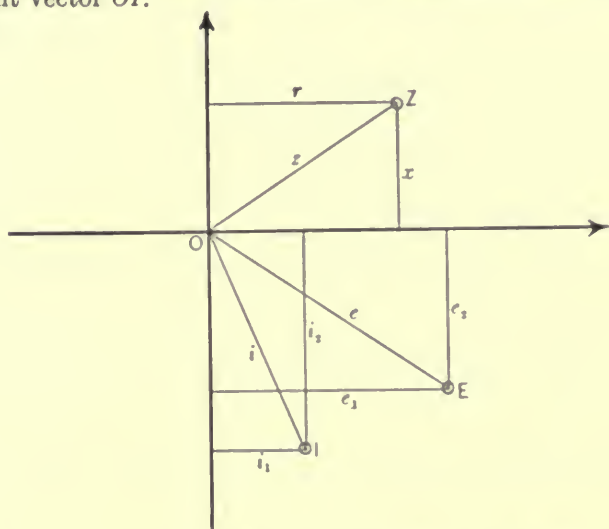


FIG. 21. Current, E.M.F. and Impedance Vector Diagram.

In the same manner an alternating E.M.F. of the same frequency can be represented by a vector $\overline{OE}$ in the same Fig. 21, and denoted by a general number,

$$E = e_1 + je_2.$$

An impedance can be represented by a general number,

$$Z = r + jx,$$

where r is the resistance and x the reactance.

If now we have two impedances, OZ_1 and OZ_2 , $Z_1 = r_1 + jx_1$ and $Z_2 = r_2 + jx_2$, their product $Z_1 Z_2$ can be formed mathematically, but it has no physical meaning.

If we have a current and a voltage, $\dot{I} = i_1 + ji_2$ and $\dot{E} = e_1 + je_2$, the product of current and voltage is the power P of the alternating circuit.

The product of the two general numbers $\dot{I}$ and $\dot{E}$ can be formed mathematically, $\dot{I}\dot{E}$, and would represent a point C in the vector plane Fig. 21. This point C , however, and the mathematical expression $\dot{I}\dot{E}$, which represents it, does not give the power P of the alternating circuit, since the power P is not of the same frequency as $\dot{I}$ and $\dot{E}$, and therefore cannot be represented in the same polar diagram Fig. 21, which represents $\dot{I}$ and $\dot{E}$.

If we have a current $\dot{I}$ and an impedance Z , in Fig. 21; $\dot{I} = i_1 + ji_2$ and $Z = r + jx$, their product is a voltage, and as the voltage is of the same frequency as the current, it can be represented in the same polar diagram, Fig. 21, and thus is given by the mathematical product of $\dot{I}$ and Z ,

$$\begin{aligned}\dot{E} &= \dot{I}Z = (i_1 + ji_2)(r + jx), \\ &= (i_1r - i_2x) + j(i_2r + i_1x).\end{aligned}$$

28. Commonly, in the denotation of graphical diagrams by general numbers, as the polar diagram of alternating currents, those quantities, which are vectors in the polar diagram, as the current, voltage, etc., are represented by dotted capitals: $\dot{E}$, $\dot{I}$, while those general numbers, as the impedance, admittance, etc., which appear as operators, that is, as multipliers of one vector, for instance the current, to get another vector, the voltage, are represented algebraically by capitals without dot: $Z = r + jx =$ impedance, etc.

This limitation of calculation with the mathematical representation of physical quantities must constantly be kept in mind in all theoretical investigations.

Division of General Numbers.

29. The division of two general numbers, $\dot{A} = a_1 + ja_2$ and $\dot{B} = b_1 + jb_2$, gives,

$$\dot{C} = \frac{\dot{A}}{\dot{B}} = \frac{a_1 + ja_2}{b_1 + jb_2}.$$

This fraction contains the quadrature number in the numerator as well as in the denominator. The quadrature number

can be eliminated from the denominator by multiplying numerator and denominator by the conjugate quantity of the denominator, $b_1 - jb_2$, which gives:

$$\begin{aligned} C &= \frac{(a_1 + ja_2)(b_1 - jb_2)}{(b_1 + jb_2)(b_1 - jb_2)} = \frac{(a_1b_1 + a_2b_2) + j(a_2b_1 - a_1b_2)}{b_1^2 + b_2^2} \\ &= \frac{a_1b_1 + a_2b_2}{b_1^2 + b_2^2} + j \frac{a_2b_1 - a_1b_2}{b_1^2 + b_2^2}; \end{aligned}$$

for instance,

$$\begin{aligned} C &= \frac{\dot{A}}{\dot{B}} = \frac{6 + 2.5j}{3 + 4j} \\ &= \frac{(6 + 2.5j)(3 - 4j)}{(3 + 4j)(3 - 4j)} \\ &= \frac{28 - 16.5j}{25} \\ &= 1.12 - 0.66j. \end{aligned}$$

If desired, the quadrature number may be eliminated from the numerator and left in the denominator by multiplying with the conjugate number of the numerator, thus:

$$\begin{aligned} C &= \frac{\dot{A}}{\dot{B}} = \frac{a_1 + ja_2}{b_1 + jb_2} \\ &= \frac{(a_1 + ja_2)(a_1 - ja_2)}{(b_1 + jb_2)(a_1 - ja_2)} \\ &= \frac{a_1^2 + a_2^2}{(a_1b_1 + a_2b_2) + j(a_1b_2 - a_2b_1)}; \end{aligned}$$

for instance,

$$\begin{aligned} C &= \frac{\dot{A}}{\dot{B}} = \frac{6 + 2.5j}{3 + 4j} \\ &= \frac{(6 + 2.5j)(6 - 2.5j)}{(3 + 4j)(6 - 2.5j)} \\ &= \frac{42.25}{28 + 16.5j} \end{aligned}$$

30. Just as in multiplication, the polar representation of the general number in division is more perspicuous than any other.

Let $\dot{A} = a(\cos \alpha + j \sin \alpha)$ be divided by $\dot{B} = b(\cos \beta + j \sin \beta)$, thus:

$$\begin{aligned} \dot{C} = \frac{\dot{A}}{\dot{B}} &= \frac{a(\cos \alpha + j \sin \alpha)}{b(\cos \beta + j \sin \beta)} \\ &= \frac{a(\cos \alpha + j \sin \alpha)(\cos \beta - j \sin \beta)}{b(\cos \beta + j \sin \beta)(\cos \beta - j \sin \beta)} \\ &= \frac{a\{(\cos \alpha \cos \beta + \sin \alpha \sin \beta) + j(\sin \alpha \cos \beta - \cos \alpha \sin \beta)\}}{b(\cos^2 \beta + \sin^2 \beta)} \\ &= \frac{a}{b}\{\cos(\alpha - \beta) + j \sin(\alpha - \beta)\}. \end{aligned}$$

That is, general numbers $\dot{A}$ and $\dot{B}$ are divided by dividing their vectors or absolute values, a and b , and subtracting their phases or angles α and β .

Involution and Evolution of General Numbers.

31. Since involution is multiple multiplication, and evolution is involution with fractional exponents, both can be resolved into simple expressions by using the polar form of the general number.

If,

$$\dot{A} = a_1 + ja_2 = a(\cos \alpha + j \sin \alpha),$$

then

$$\dot{C} = \dot{A}^n = a^n(\cos n\alpha + j \sin n\alpha).$$

For instance, if

$$\dot{A} = 3 + 4j = 5(\cos 53 \text{ deg.} + j \sin 53 \text{ deg.});$$

then,

$$\begin{aligned} \dot{C} = \dot{A}^4 &= 5^4(\cos 4 \times 53 \text{ deg.} + j \sin 4 \times 53 \text{ deg.}) \\ &= 625(\cos 212 \text{ deg.} + j \sin 212 \text{ deg.}) \\ &= 625(-\cos 32 \text{ deg.} - j \sin 32 \text{ deg.}) \\ &= 625(-0.848 - 0.530 j) \\ &= -529 - 331 j. \end{aligned}$$

If, $\dot{A} = a_1 + ja_2 = a(\cos \alpha + j \sin \alpha)$, then

$$\begin{aligned} \dot{C} = \sqrt[n]{\dot{A}} &= \dot{A}^{\frac{1}{n}} = a^{\frac{1}{n}} \left(\cos \frac{\alpha}{n} + j \sin \frac{\alpha}{n} \right) \\ &= \sqrt[n]{a} \left(\cos \frac{\alpha}{n} + j \sin \frac{\alpha}{n} \right). \end{aligned}$$

32. If, in the polar expression of A , we increase the phase angle α by 2π , or by any multiple of 2π : $2q\pi$, where q is any integer number, we get the same value of A , thus:

$$A = a\{\cos(\alpha + 2q\pi) + j \sin(\alpha + 2q\pi)\},$$

since the cosine and sine repeat after every 360 deg, or 2π .

The n th root, however, is different:

$$C = \sqrt[n]{A} = \sqrt[n]{a} \left(\cos \frac{\alpha + 2q\pi}{n} + j \sin \frac{\alpha + 2q\pi}{n} \right).$$

We hereby get n different values of C , for $q=0, 1, 2, \dots, n-1$; $q=n$ gives again the same as $q=0$. Since it gives

$$\frac{\alpha + 2n\pi}{n} = \frac{\alpha}{n} + 2\pi;$$

that is, an increase of the phase angle by 360 deg., which leaves cosine and sine unchanged.

Thus, the n th root of any general number has n different values, and these values have the same vector or absolute term $\sqrt[n]{a}$, but differ from each other by the phase angle $\frac{2\pi}{n}$ and its multiples.

For instance, let $A = -529 - 331j = 625 (\cos 212 \text{ deg.} + j \sin 212 \text{ deg.})$ then,

$$\begin{aligned} C = \sqrt[n]{A} &= \sqrt[4]{625} \left(\cos \frac{212 + 360q}{4} + j \sin \frac{212 + 360q}{4} \right) \\ &= 5(\cos 53 + j \sin 53) &= 3 + 4j \\ &= 5(\cos 143 + j \sin 143) = 5(-\cos 37 + j \sin 37) = -4 + 3j \\ &= 5(\cos 233 + j \sin 233) = 5(-\cos 53 - j \sin 53) = -3 - 4j \\ &= 5(\cos 323 + j \sin 323) = 5(\cos 37 - j \sin 37) = 4 - 3j \\ &= 5(\cos 413 + j \sin 413) = 5(\cos 53 + j \sin 53) = 3 + 4j \end{aligned}$$

The n roots of a general number $A = a(\cos \alpha + j \sin \alpha)$ differ from each other by the phase angles $\frac{2\pi}{n}$, or 1/nth of 360 deg., and since they have the same absolute value $\sqrt[n]{a}$, it follows, that they are represented by n equidistant points of a circle with radius $\sqrt[n]{a}$, as shown in Fig. 22, for $n=4$, and in Fig. 23 for

$n=9$. Such a system of n equal vectors, differing in phase from each other by $1/n$ th of 360 deg., is called a *polyphase system*, or an *n -phase system*. The n roots of the general number thus give an n -phase system.

33. For instance, $\sqrt[n]{1}=?$

If $A=a(\cos \alpha + j \sin \alpha)=1$, this means: $a=1$, $\alpha=0$; and hence,

$$\sqrt[n]{1} = \cos \frac{2q\pi}{n} + j \sin \frac{2q\pi}{n};$$

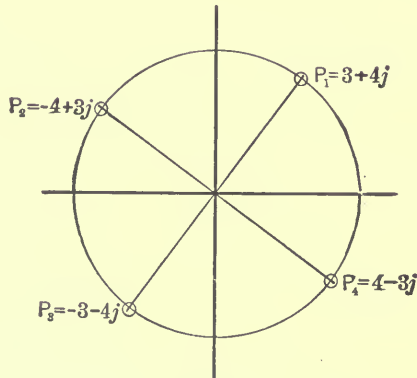


FIG. 22. Roots of a General Number, $n=4$.

and the n roots of the unit are

$$q=0 \quad \sqrt[n]{1}=1;$$

$$q=1 \quad \cos \frac{360}{n} + j \sin \frac{360}{n};$$

$$q=2 \quad \cos 2 \times \frac{360}{n} + j \sin 2 \times \frac{360}{n};$$

$$q=n-1 \quad \cos (n-1) \frac{360}{n} + j \sin (n-1) \frac{360}{n}.$$

However,

$$\cos q \frac{360}{n} + j \sin q \frac{360}{n} = \left(\cos \frac{360}{n} + j \sin \frac{360}{n} \right)^q;$$

hence, the n roots of 1 are,

$$\sqrt[n]{1} = \left(\cos \frac{360}{n} + j \sin \frac{360}{n} \right)^q,$$

where q may be any integer number.

One of these roots is real, for $q=0$, and is $= +1$.

If n is odd, all the other roots are general, or complex numbers.

If n is an even number, a second root, for $q = \frac{n}{2}$, is also real: $\cos 180 + j \sin 180 = -1$.

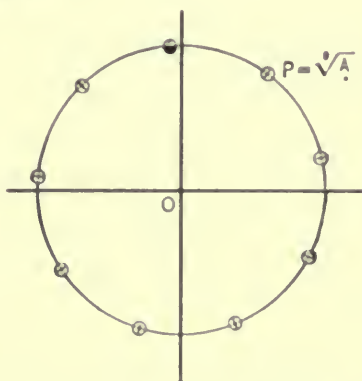


FIG. 23. Roots of a General Number, $n=9$.

If n is divisible by 4, two roots are quadrature numbers, and are $+j$, for $q = \frac{n}{4}$, and $-j$, for $q = \frac{3n}{4}$.

34. Using the rectangular coordinate expression of the general number, $A = a_1 + ja_2$, the calculation of the roots becomes more complicated. For instance, given $\sqrt[4]{A} = ?$

Let $C = \sqrt[4]{A} = c_1 + jc_2$;
then, squaring,

$$A = (c_1 + jc_2)^2;$$

hence,

$$a_1 + ja_2 = (c_1^2 - c_2^2) + 2jc_1c_2.$$

Since, if two general numbers are equal, their horizontal and their vertical components must be equal, it is:

$$a_1 = c_1^2 - c_2^2 \quad \text{and} \quad a_2 = 2c_1c_2.$$

Squaring both equations and adding them, gives,

$$a_1^2 + a_2^2 = (c_1^2 + c_2^2)^2.$$

Hence:

$$c_1^2 + c_2^2 = \sqrt{a_1^2 + a_2^2},$$

and since

$$c_1^2 - c_2^2 = a_1;$$

then,

$$c_1^2 = \frac{1}{2}(\sqrt{a_1^2 + a_2^2} + a_1),$$

and

$$c_2^2 = \frac{1}{2}(\sqrt{a_1^2 + a_2^2} - a_1).$$

Thus

$$c_1 = \sqrt{\frac{1}{2}\{\sqrt{a_1^2 + a_2^2} + a_1\}}$$

and

$$c_2 = \sqrt{\frac{1}{2}\{\sqrt{a_1^2 + a_2^2} - a_1\}},$$

and

$$\sqrt[3]{A} = \sqrt{\frac{1}{2}\{\sqrt{a_1^2 + a_2^2} + a_1\}} + j\sqrt{\frac{1}{2}\{\sqrt{a_1^2 + a_2^2} - a_1\}},$$

which is a rather complicated expression.

35. When representing physical quantities by general numbers, that is, complex quantities, at the end of the calculation the final result usually appears also as a general number, or as a complex of general numbers, and then has to be reduced to the absolute value and the phase angle of the physical quantity. This is most conveniently done by reducing the general numbers to their polar expressions. For instance, if the result of the calculation appears in the form,

$$R = \frac{(a_1 + ja_2)(b_1 + jb_2)^3\sqrt{c_1 + jc_2}}{(d_1 + jd_2)^2(e_1 + je_2)},$$

by substituting

$$a = \sqrt{a_1^2 + a_2^2}; \quad \tan \alpha = \frac{a_2}{a_1}.$$

$$b = \sqrt{b_1^2 + b_2^2}; \quad \tan \beta = \frac{b_2}{b_1};$$

and so on.

$$\begin{aligned} R &= \frac{a(\cos \alpha + j \sin \alpha)b^3(\cos \beta + j \sin \beta)^3\sqrt{c}(\cos \gamma + j \sin \gamma)^{\frac{1}{2}}}{d^2(\cos \delta + j \sin \delta)^2e(\cos \epsilon + j \sin \epsilon)} \\ &= \frac{ab^3\sqrt{c}}{d^2e} \{ \cos(\alpha + 3\beta + \gamma/2 - 2\delta - \epsilon) + j \sin(\alpha + 3\beta + \gamma/2 - 2\delta - \epsilon) \}. \end{aligned}$$

Therefore, the absolute value of a fractional expression is the product of the absolute values of the factors of the numerator, divided by the product of the absolute values of the factors of the denominator.

The phase angle of a fractional expression is the sum of the phase angles of the factors of the numerator, minus the sum of the phase angles of the factors of the denominator.

For instance,

$$\begin{aligned}
 R &= \frac{(3-4j)^2(2+2j)^2\sqrt{-2.5+6j}}{5(4+3j)^2\sqrt{2}} \\
 &= \frac{25(\cos 307 + j \sin 307)^2 2\sqrt{2}(\cos 45 + j \sin 45)^2 \sqrt{6.5}(\cos 114 + j \sin 114)^{\frac{1}{2}}}{125(\cos 37 + j \sin 37)^2 \sqrt{2}} \\
 &= 0.4\sqrt[3]{6.5} \left\{ \cos \left(2 \times 307 + 45 + \frac{114}{3} - 2 \times 37 \right) \right. \\
 &\quad \left. + j \sin \left(2 \times 307 + 45 + \frac{114}{3} - 2 \times 37 \right) \right\} \\
 &= 0.4\sqrt[3]{6.5} \{ \cos 263 + j \sin 263 \} \\
 &= 0.746 \{ -0.122 - 0.992j \} = -0.091 - 0.74j.
 \end{aligned}$$

36. As will be seen in Chapter II:

$$\begin{aligned}
 e^u &= 1 + u + \frac{u^2}{2} + \frac{u^3}{3} + \frac{u^4}{4} + \dots \\
 \cos x &= 1 - \frac{x^2}{2} + \frac{x^4}{4} - \frac{x^6}{6} + \frac{x^8}{8} - \dots \\
 \sin x &= x - \frac{x^3}{3} + \frac{x^5}{5} - \frac{x^7}{7} + \dots
 \end{aligned}$$

Herefrom follows, by substituting, $x = \theta$, $u = j\theta$,

$$\cos \theta + j \sin \theta = e^{j\theta},$$

and the polar expression of the complex quantity,

$$A = a(\cos \alpha + j \sin \alpha),$$

thus can also be written in the form,

$$A = a e^{j\alpha},$$

where ε is the base of the natural logarithms,

$$\varepsilon = 1 + 1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} + \dots = 2.71828 \dots$$

Since any number a can be expressed as a power of any other number, one can substitute,

$$a = \varepsilon^{\alpha_0},$$

where $\alpha_0 = \log_{\varepsilon} a = \frac{\log_{10} a}{\log_{10} \varepsilon}$, and the general number thus can also be written in the form,

$$A = \varepsilon^{\alpha_0 + j\alpha};$$

that is the general number, or complex quantity, can be expressed in the forms,

$$\begin{aligned} A &= a_1 + ja_2 \\ &= a(\cos \alpha + j \sin \alpha) \\ &= a\varepsilon^{j\alpha} = \varepsilon^{\alpha_0 + j\alpha}. \end{aligned}$$

The last two, or exponential forms, are rarely used, as they are less convenient for algebraic operations. They are of importance, however, since solutions of differential equations frequently appear in this form, and then are reduced to the polar or the rectangular form.

37. For instance, the differential equation of the distribution of alternating current in a flat conductor, or of alternating magnetic flux in a flat sheet of iron, has the form:

$$\frac{d^2 y}{dx^2} = -2jc^2 y;$$

and is integrated by $y = A\varepsilon^{-Vx}$, where,

$$V = \sqrt{-2jc^2} = \pm(1-j)c;$$

hence,

$$y = A_1 \varepsilon^{+(1-j)cx} + A_2 \varepsilon^{-(1-j)cx}.$$

This expression, reduced to the polar form, is

$$y = A_1 \varepsilon^{+cx} (\cos cx - j \sin cx) + A_2 \varepsilon^{-cx} (\cos cx + j \sin cx).$$

Logarithmation.

38. In taking the logarithm of a general number, the exponential expression is most convenient, thus:

$$\begin{aligned}\log_{\epsilon} (a_1 + ja_2) &= \log_{\epsilon} a (\cos \alpha + j \sin \alpha) \\ &= \log_{\epsilon} a \epsilon^{j\alpha} \\ &= \log_{\epsilon} a + \log_{\epsilon} \epsilon^{j\alpha} \\ &= \log_{\epsilon} a + j\alpha;\end{aligned}$$

or, if b = base of the logarithm, for instance, $b = 10$, it is:

$$\log_b (a_1 + ja_2) = \log_b a \epsilon^{j\alpha} = \log_b a + j\alpha \log_b \epsilon;$$

or, if b unequal 10, reduced to $\log_{10}$:

$$\log_b (a_1 + ja_2) = \frac{\log_{10} a}{\log_{10} b} + j\alpha \frac{\log_{10} \epsilon}{\log_{10} b}.$$

NOTE. In mathematics, for quadrature unit $\sqrt{-1}$ is always chosen the symbol i . Since, however, in engineering the symbol i is universally used to represent *electric current*, for the quadrature unit the symbol j has been chosen, as the letter nearest in appearance to i , and j thus is always used in engineering calculations to denote the quadrature unit $\sqrt{-1}$.

CHAPTER II.

POTENTIAL SERIES AND EXPONENTIAL FUNCTION.

A. GENERAL.

39. An expression such as

$$y = \frac{1}{1-x} \quad . \quad . \quad . \quad . \quad . \quad . \quad (1)$$

represents a fraction; that is, the result of division, and like any fraction it can be calculated; that is, the fractional form eliminated, by dividing the numerator by the denominator, thus:

$$1-x \overline{)1 = 1 + x + x^2 + x^3 + \dots}$$

$$\begin{array}{r} 1-x \\ +x \\ \hline x-x^2 \\ +x^2 \\ \hline x^2-x^3 \\ +x^3 \\ \hline \end{array}$$

Hence, the fraction (1) can also be expressed in the form:

$$y = \frac{1}{1-x} = 1 + x + x^2 + x^3 + \dots \quad . \quad . \quad . \quad (2)$$

This is an infinite series of successive powers of x , or a *potential series*.

In the same manner, by dividing through, the expression

$$y = \frac{1}{1+x}, \quad . \quad . \quad . \quad . \quad . \quad . \quad (3)$$

can be reduced to the infinite series,

$$y = \frac{1}{1+x} = 1 - x + x^2 - x^3 + \dots \quad . \quad . \quad . \quad (4)$$

The infinite series (2) or (4) is another form of representation of the expression (1) or (3), just as the periodic decimal fraction is another representation of the common fraction (for instance $0.6363 \dots = 7/11$).

40. As the series contains an infinite number of terms, in calculating numerical values from such a series perfect exactness can never be reached; since only a finite number of terms are calculated, the result can only be an approximation. By taking a sufficient number of terms of the series, however, the approximation can be made as close as desired; that is, numerical values may be calculated as exactly as necessary, so that for engineering purposes the infinite series (2) or (4) gives just as exact numerical values as calculation by a finite expression (1) or (2), provided a sufficient number of terms are used. In most engineering calculations, an exactness of 0.1 per cent is sufficient; rarely is an exactness of 0.01 per cent or even greater required, as the unavoidable variations in the nature of the materials used in engineering structures, and the accuracy of the measuring instruments impose a limit on the exactness of the result.

For the value $x=0.5$, the expression (1) gives $y = \frac{1}{1-0.5} = 2$; while its representation by the series (2) gives

$$y = 1 + 0.5 + 0.25 + 0.125 + 0.0625 + 0.03125 + \dots \quad (5)$$

and the successive approximations of the numerical values of y then are:

using one term:	$y = 1$	$= 1$;	error: -1
" two terms:	$y = 1 + 0.5$	$= 1.5$;	" -0.5
" three terms:	$y = 1 + 0.5 + 0.25$	$= 1.75$;	" -0.25
" four terms:	$y = 1 + 0.5 + 0.25 + 0.125$	$= 1.875$;	" -0.125
" five terms:	$y = 1 + 0.5 + 0.25 + 0.125 + 0.0625$	$= 1.9375$	" -0.0625

It is seen that the successive approximations come closer and closer to the correct value, $y=2$, but in this case always remain below it; that is, the series (2) approaches its limit from below, as shown in Fig. 24, in which the successive approximations are marked by crosses.

For the value $x=0.5$, the approach of the successive approximations to the limit is rather slow, and to get an accuracy of 0.1 per cent, that is, bring the error down to less than 0.002, requires a considerable number of terms.

For $x=0.1$ the series (2) is

$$y = 1 + 0.1 + 0.01 + 0.001 + 0.0001 + \dots \quad (6)$$

and the successive approximations thus are

$$1:y=1; \quad 2:y=1.1; \quad 3:y=1.11; \quad 4:y=1.111; \quad 5:y=1.1111;$$

and as, by (1), the final or limiting value is

$$y = \frac{1}{1-0.1} = \frac{10}{9} = 1.1111 \dots$$

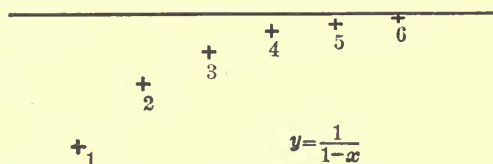


FIG. 24. Convergent Series with One-sided Approach.

the fourth approximation already brings the error well below 0.1 per cent, and sufficient accuracy thus is reached for most engineering purposes by using four terms of the series.

41. The expression (3) gives, for $x=0.5$, the value,

$$y = \frac{1}{1+0.5} = \frac{2}{3} = 0.6666 \dots$$

Represented by series (4), it gives

$$y = 1 - 0.5 + 0.25 - 0.125 + 0.0625 - 0.03125 + \dots \quad (7)$$

the successive approximations are;

1st: $y=1$	$=1;$	error: $+0.333 \dots$
2d: $y=1-0.5$	$=0.5;$	" $-0.1666 \dots$
3d: $y=1-0.5+0.25$	$=0.75;$	" $+0.0833 \dots$
4th: $y=1-0.5+0.25-0.125$	$=0.625;$	" $-0.04166 \dots$
5th: $y=1-0.5+0.25-0.125+0.0625$	$=0.6875;$	" $+0.020833 \dots$

As seen, the successive approximations of this series come closer and closer to the correct value $y=0.6666 \dots$, but in this case are alternately above and below the correct or limiting value, that is, the series (4) approaches its limit from both sides, as shown in Fig. 25, while the series (2) approached the limit from below, and still other series may approach their limit from above.

With such alternating approach of the series to the limit, as exhibited by series (4), the limiting or final value is between any two successive approximations, that is, the error of any approximation is less than the difference between this and the next following approximation.

Such a series thus is preferable in engineering, as it gives information on the maximum possible error, while the series with one-sided approach does not do this without special investigation, as the error is greater than the difference between successive approximations.

42. Substituting $x=2$ into the expressions (1) and (2), equation (1) gives

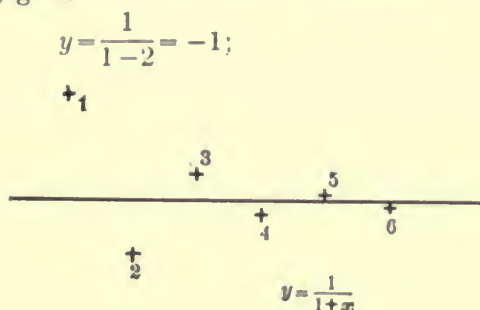


FIG. 25. Convergent Series with Alternating Approach.

while the infinite series (2) gives

$$y = 1 + 2 + 4 + 8 + 16 + 32 + \dots;$$

and the successive approximations of the latter thus are

$$1; \quad 3; \quad 7; \quad 15; \quad 31; \quad 63. \dots;$$

that is, the successive approximations do not approach closer and closer to a final value, but, on the contrary, get further and further away from each other, and give entirely wrong results. They give increasing positive values, which apparently approach ∞ for the entire series, while the correct value of the expression, by (1), is $y = -1$.

Therefore, for $x=2$, the series (2) gives unreasonable results, and thus cannot be used for calculating numerical values.

The same is the case with the representation (4) of the expression (3) for $x=2$. The expression (3) gives

$$y = \frac{1}{1+2} = 0.3333 \dots;$$

while the infinite series (4) gives

$$y = 1 - 2 + 4 - 8 + 16 - 32 + \dots,$$

and the successive approximations of the latter thus are

$$1; \quad -1; \quad +3; \quad -5; \quad +11; \quad -21; \dots;$$

hence, while the successive values still are alternately above and below the correct or limiting value, they do not approach it with increasing closeness, but more and more diverge therefrom.

Such a series, in which the values derived by the calculation of more and more terms do not approach a final value closer and closer, is called *divergent*, while a series is called *convergent* if the successive approximations approach a final value with increasing closeness.

43. While a finite expression, as (1) or (3), holds good for all values of x , and numerical values of it can be calculated whatever may be the value of the independent variable x , an infinite series, as (2) and (4), frequently does not give a finite result for every value of x , but only for values within a certain range. For instance, in the above series, for $-1 < x < +1$, the series is convergent; while for values of x outside of this range the series is divergent and thus useless.

When representing an expression by an infinite series, it thus is necessary to determine that the series is convergent; that is, approaches with increasing number of terms a finite limiting value, otherwise the series cannot be used. Where the series is convergent within a certain range of x , divergent outside of this range, it can be used only in the *range of convergency*, but outside of this range it cannot be used for deriving numerical values, but some other form of representation has to be found which is convergent.

This can frequently be done, and the expression thus represented by one series in one range and by another series in another range. For instance, the expression (1), $y = \frac{1}{1+x}$, by substituting, $x = \frac{1}{u}$, can be written in the form

$$y = \frac{1}{1 + \frac{1}{u}} = + \frac{u}{1+u},$$

and then developed into a series by dividing the numerator by the denominator, which gives

$$y = u - u^2 + u^3 - u^4 + \dots;$$

or, resubstituting x ,

$$y = \frac{1}{x} - \frac{1}{x^2} + \frac{1}{x^3} - \frac{1}{x^4} + \dots \quad (8)$$

which is convergent for $x=2$, and for $x=2$ it gives

$$y = 0.5 - 0.25 + 0.125 - 0.0625 + \dots \quad (9)$$

With the successive approximations:

$$0.5; \quad 0.25; \quad 0.375; \quad 0.3125, \dots$$

which approach the final limiting value,

$$y = 0.333 \dots$$

44. An infinite series can be used only if it is convergent. Mathematical methods exist for determining whether a series is convergent or not. For engineering purposes, however, these methods usually are unnecessary: for practical use it is not sufficient that a series be convergent, but it must converge so rapidly—that is, the successive terms of the series must decrease at such a great rate—that accurate numerical results are derived by the calculation of only a very few terms; two or three, or perhaps three or four. This, for instance, is the case with the series (2) and (4) for $x=0.1$ or less. For $x=0.5$, the series (2) and (4) are still convergent, as seen in (5) and (7), but are useless for most engineering purposes, as the successive terms decrease so slowly that a large number of terms have to be calculated to get accurate results, and for such lengthy calculations there is no time in engineering work. If, however, the successive terms of a series decrease at such a rapid rate that all but the first few terms can be neglected, the series is certain to be convergent.

In a series therefore, in which there is a question whether it is convergent or divergent, as for instance the series

$$y = 1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} + \frac{1}{5} + \frac{1}{6} + \dots \text{ (divergent),}$$

or

$$y = 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \frac{1}{5} - \frac{1}{6} + \dots \text{ (convergent),}$$

the matter of convergency is of little importance for engineering calculation, as the series is useless in any case; that is, does not give accurate numerical results with a reasonably moderate amount of calculation.

A series, to be usable for engineering work, must have the successive terms decreasing at a very rapid rate, and if this is the case, the series is convergent, and the mathematical investigations of convergency thus usually becomes unnecessary in engineering work.

45. It would rarely be advantageous to develop such simple expressions as (1) and (3) into infinite series, such as (2) and (4), since the calculation of numerical values from (1) and (3) is simpler than from the series (2) and (4), even though very few terms of the series need to be used.

The use of the series (2) or (4) instead of the expressions (1) and (3) therefore is advantageous only if these series converge so rapidly that only the first two terms are required for numerical calculation, and the third term is negligible; that is, for very small values of x . Thus, for $x=0.01$, according to (2),

$$y = 1 + 0.01 + 0.0001 + \dots = 1 + 0.01,$$

as the next term, 0.0001, is already less than 0.01 per cent of the value of the total expression.

For very small values of x , therefore, by (1) and (2),

$$y = \frac{1}{1-x} = 1 + x, \quad \dots \dots \dots (10)$$

and by (3) and (4),

$$y = \frac{1}{1+x} = 1 - x, \quad \dots \dots \dots (11)$$

and these expressions (10) and (11) are useful and very commonly used in engineering calculation for simplifying work. For instance, if 1 plus or minus a very small quantity appears as factor in the denominator of an expression, it can be replaced by 1 minus or plus the same small quantity as factor in the numerator of the expression, and inversely.

For example, if a direct-current receiving circuit, of resistance r , is fed by a supply voltage e_0 over a line of low resistance r_0 , what is the voltage e at the receiving circuit?

The total resistance is $r + r_0$; hence, the current, $i = \frac{e_0}{r + r_0}$, and the voltage at the receiving circuit is

$$e = ri = e_0 \frac{r}{r + r_0}. \quad (12)$$

If now r_0 is small compared with r , it is

$$e = e_0 \frac{1}{1 + \frac{r_0}{r}} = e_0 \left\{ 1 - \frac{r_0}{r} \right\} \quad (13)$$

As the next term of the series would be $\left(\frac{r_0}{r}\right)^2$, the error made by the simpler expression (13) is less than $\left(\frac{r_0}{r}\right)^2$. Thus, if r_0 is 3 per cent of r , which is a fair average in interior lighting circuits, $\left(\frac{r_0}{r}\right)^2 = 0.03^2 = 0.0009$, or less than 0.1 per cent; hence, is usually negligible.

46. If an expression in its finite form is more complicated and thereby less convenient for numerical calculation, as for instance if it contains roots, development into an infinite series frequently simplifies the calculation.

Very convenient for development into an infinite series of powers or roots, is the *binomial theorem*,

$$(1 \pm u)^n = 1 \pm nu + \frac{n(n-1)}{2} u^2 \pm \frac{n(n-1)(n-2)}{3} u^3 + \dots \left. \vphantom{\frac{n(n-1)(n-2)}{3}} \right\} \quad (14)$$

where $\underline{m} = 1 \times 2 \times 3 \times \dots \times m$.

Thus, for instance, in an alternating-current circuit of resistance r , reactance x , and supply voltage e , the current is,

$$i = \frac{e}{\sqrt{r^2 + x^2}} \quad (15)$$

If this circuit is practically non-inductive, as an incandescent lighting circuit; that is, if x is small compared with r , (15) can be written in the form,

$$i = \frac{e}{r\sqrt{1 + \left(\frac{x}{r}\right)^2}} = \frac{e}{r} \left[1 + \left(\frac{x}{r}\right)^2 \right]^{-\frac{1}{2}}, \quad \dots \quad (16)$$

and the square root can be developed by the binomial (14), thus, $u = \left(\frac{x}{r}\right)^2$; $n = -\frac{1}{2}$, and gives

$$\left[1 + \left(\frac{x}{r}\right)^2 \right]^{-\frac{1}{2}} = 1 - \frac{1}{2} \left(\frac{x}{r}\right)^2 + \frac{3}{8} \left(\frac{x}{r}\right)^4 - \frac{5}{16} \left(\frac{x}{r}\right)^6 + \dots \quad (17)$$

In this series (17), if $x = 0.1r$ or less; that is, the reactance is not more than 10 per cent of the resistance, the third term, $\frac{3}{8} \left(\frac{x}{r}\right)^4$, is less than 0.01 per cent; hence, negligible, and the series is approximated with sufficient exactness by the first two terms,

$$\left[1 + \left(\frac{x}{r}\right)^2 \right]^{-\frac{1}{2}} = 1 - \frac{1}{2} \left(\frac{x}{r}\right)^2, \quad \dots \quad (18).$$

and equation (16) of the current then gives

$$i = \frac{e}{r} \left\{ 1 - \frac{1}{2} \left(\frac{x}{r}\right)^2 \right\}. \quad \dots \quad (19)$$

This expression is simpler for numerical calculations than the expression (15), as it contains no square root.

47. Development into a series may become necessary, if further operations have to be carried out with an expression for which the expression is not suited, or at least not well suited. This is often the case where the expression has to be integrated, since very few expressions can be integrated.

Expressions under an integral sign therefore very commonly have to be developed into an infinite series to carry out the integration.

EXAMPLE 1.

Of the equilateral hyperbola (Fig. 26),

$$xy = a^2, \quad (20)$$

the length L of the arc between $x_1 = 2a$ and $x_2 = 10a$ is to be calculated.

An element dl of the arc is the hypotenuse of a right triangle with dx and dy as cathetes. It, therefore, is,

$$\begin{aligned} dl &= \sqrt{dx^2 + dy^2} \\ &= \sqrt{1 + \left(\frac{dy}{dx}\right)^2} dx, \quad (21) \end{aligned}$$

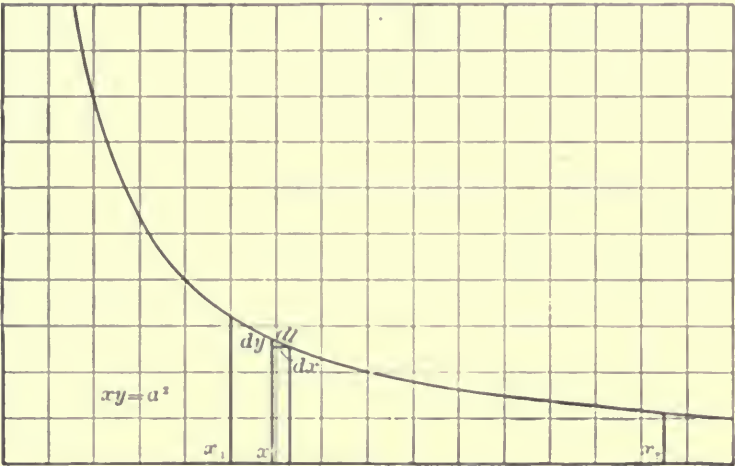


FIG. 26. Equilateral Hyperbola.

and from (20),

$$y = \frac{a^2}{x} \quad \text{and} \quad \frac{dy}{dx} = -\frac{a^2}{x^2}, \quad (22)$$

Substituting (22) in (21) gives,

$$dl = \sqrt{1 + \left(\frac{a}{x}\right)^4} dx; \quad (23)$$

hence, the length L of the arc, from x_1 to x_2 is,

$$L = \int_{x_1}^{x_2} dl = \int_{x_1}^{x_2} \sqrt{1 + \left(\frac{a}{x}\right)^4} dx.$$

Substituting $\frac{x}{a} = v$; that is, $dx = a dv$, also substituting

$$v_1 = \frac{x_1}{a} = 2 \quad \text{and} \quad v_2 = \frac{x_2}{a} = 10,$$

gives

$$L = a \int_{v_1}^{v_2} \sqrt{1 + \frac{1}{v^4}} dv.$$

The expression under the integral is inconvenient for integration; it is preferably developed into an infinite series, by the binomial theorem (14).

Write $u = \frac{1}{v^4}$ and $n = \frac{1}{2}$, then

$$\sqrt{1 + \frac{1}{v^4}} = 1 + \frac{1}{2v^4} - \frac{1}{8v^8} + \frac{1}{16v^{12}} - \frac{5}{128v^{16}} + \dots,$$

and

$$\begin{aligned} L &= a \int_{v_1}^{v_2} \left\{ 1 + \frac{1}{2v^4} - \frac{1}{8v^8} + \frac{1}{16v^{12}} - \frac{5}{128v^{16}} + \dots \right\} dv \\ &= av \left\{ 1 - \frac{1}{2 \times 3 \times v^4} + \frac{1}{7 \times 8 \times v^8} - \frac{1}{11 \times 16 \times v^{12}} \right. \\ &\quad \left. + \frac{1}{3 \times 128 \times v^{16}} - \dots \right\}_{v_1}^{v_2} \\ &= a \left\{ (v_2 - v_1) + \frac{1}{6} \left(\frac{1}{v_1^3} - \frac{1}{v_2^3} \right) - \frac{1}{56} \left(\frac{1}{v_1^7} - \frac{1}{v_2^7} \right) \right. \\ &\quad \left. + \frac{1}{176} \left(\frac{1}{v_1^{11}} - \frac{1}{v_2^{11}} \right) - \dots \right\}, \end{aligned}$$

and substituting the numerical values,

$$\begin{aligned} L &= a \left\{ (10 - 2) + \frac{1}{6}(0.125 - 0.001) \right. \\ &\quad \left. - \frac{1}{56}(0.0078 - 0) + \frac{1}{176}(0.0001 - 0) \right\} \\ &= a \{ 8 + 0.0207 - 0.0001 \} = 8.0206a. \end{aligned}$$

As seen, in this series, only the first two terms are appreciable in value, the third term less than 0.01 per cent of the total, and hence negligible, therefore the series converges very rapidly, and numerical values can easily be calculated by it.

For $x_1 < 2a$; that is, $v_1 < 2$, the series converges less rapidly, and becomes divergent for $x_1 < a$; or, $v_1 < 1$. Thus this series (17) is convergent for $v > 1$, but near this limit of convergency it is of no use for engineering calculation, as it does not converge with sufficient rapidity, and it becomes suitable for engineering calculation only when v_1 approaches 2.

EXAMPLE 2.

48. $\log 1 = 0$, and, therefore $\log (1+x)$ is a small quantity if x is small. $\log (1+x)$ shall therefore be developed in such a series of powers of x , which permits its rapid calculation without using logarithm tables.

It is

$$\log u = \int \frac{du}{u};$$

then, substituting $(1+x)$ for u gives,

$$\log (1+x) = \int \frac{dx}{1+x}. \quad \dots \dots \dots (24)$$

From equation (4)

$$\frac{1}{1+x} = 1 - x + x^2 - x^3 + \dots$$

hence, substituted into (24),

$$\begin{aligned} \log (1+x) &= \int (1 - x + x^2 - x^3 + \dots) dx \\ &= \int dx - \int x dx + \int x^2 dx - \int x^3 dx + \dots \\ &= x - \frac{x^2}{2} + \frac{x^3}{3} - \frac{x^4}{4} + \dots \dots \dots (25) \end{aligned}$$

hence, if x is very small, $\frac{x^2}{2}$ is negligible, and, therefore, all terms beyond the first are negligible, thus,

$$\log (1+x) = x;$$

while, if the second term is still appreciable in value, the more complete, but still fairly simple expression can be used,

$$\log (1+x) = x - \frac{x^2}{2} = x \left(1 - \frac{x}{2} \right)$$

If instead of the natural logarithm, as used above, the decimal logarithm is required, the following relation may be applied:

$$\log_{10} a = \log_{10} e \log_e a = 0.4343 \log_e a,$$

$\log_{10} a$ is expressed by $\log_e a$, and thus (19), (20) (21) assume the form,

$$\log_{10} (1+x) = 0.4343 \left(x - \frac{x^2}{2} + \frac{x^3}{3} - \frac{x^4}{4} + \dots \right);$$

or, approximately,

$$\log_{10}(1+x) = 0.4343x;$$

or, more accurately,

$$\log_{10} (1+x) = 0.4343x \left(1 - \frac{x}{2} \right)$$

B. DIFFERENTIAL EQUATIONS.

49. The representation by an infinite series is of special value in those cases, in which no finite expression of the function is known, as for instance, if the relation between x and y is given by a differential equation.

Differential equations are solved by separating the variables, that is, bringing the terms containing the one variable, y , on one side of the equation, the terms with the other variable x on the other side of the equation, and then separately integrating both sides of the equation. Very rarely, however, is it possible to separate the variables in this manner, and where it cannot be done, usually no systematic method of solving the differential equation exists, but this has to be done by trying different functions, until one is found which satisfies the equation.

In electrical engineering, currents and voltages are dealt with as functions of time. The current and e.m.f. giving the power lost in resistance are related to each other by Ohm's law. Current also produces a magnetic field, and this magnetic field by its changes generates an e.m.f.—the e.m.f. of self-inductance. In this case, e.m.f. is related to the change of current; that is, the differential coefficient of the current, and thus also to the differential coefficient of e.m.f., since the e.m.f.

is related to the current by Ohm's law. In a condenser, the current and therefore, by Ohm's law, the e.m.f., depends upon and is proportional to the rate of change of the e.m.f. impressed upon the condenser; that is, it is proportional to the differential coefficient of e.m.f.

Therefore, in circuits having resistance and inductance, or resistance and capacity, a relation exists between currents and e.m.fs., and their differential coefficients, and in circuits having resistance, inductance and capacity, a double relation of this kind exists; that is, a relation between current or e.m.f. and their first and second differential coefficients.

The most common differential equations of electrical engineering thus are the relations between the function and its differential coefficient, which in its simplest form is,

$$\frac{dy}{dx} = y; \quad \dots \dots \dots (26)$$

or

$$\frac{dy}{dx} = ay, \quad \dots \dots \dots (27)$$

and where the circuit has capacity as well as inductance, the second differential coefficient also enters, and the relation in its simplest form is,

$$\frac{d^2y}{dx^2} = y; \quad \dots \dots \dots (28)$$

or

$$\frac{d^2y}{dx^2} = ay, \quad \dots \dots \dots (29)$$

and the most general form of this most common differential equation of electrical engineering then is,

$$\frac{d^2y}{dx^2} + 2c \frac{dy}{dx} + ay + b = 0. \quad \dots \dots \dots (30)$$

The differential equations (26) and (27) can easily be integrated by separating the variables, but not so with equations (28), (29) and (30); the latter are preferably solved by trial.

50. The general method of solution may be illustrated with the equation (26):

$$\frac{dy}{dx} = y. \quad \dots \dots \dots (26)$$

To determine whether this equation can be integrated by an infinite series, choose such an infinite series, and then, by substituting it into equation (26), ascertain whether it satisfies the equation (26); that is, makes the left side equal to the right side for every value of x .

Let,

$$y = a_0 + a_1x + a_2x^2 + a_3x^3 + a_4x^4 + \dots \quad (31)$$

be an infinite series, of which the coefficients $a_0, a_1, a_2, a_3, \dots$ are still unknown, and by substituting (31) into the differential equation (26), determine whether such values of these coefficients can be found, which make the series (31) satisfy the equation (26).

Differentiating (31) gives,

$$\frac{dy}{dx} = a_1 + 2a_2x + 3a_3x^2 + 4a_4x^3 + \dots \quad (32)$$

The differential equation (26) transposed gives,

$$\frac{dy}{dx} - y = 0. \quad (33)$$

Substituting (31) and (32) into (33), and arranging the terms in the order of x , gives,

$$(a_1 - a_0) + (2a_2 - a_1)x + (3a_3 - a_2)x^2 + (4a_4 - a_3)x^3 + (5a_5 - a_4)x^4 + \dots = 0. \quad (34)$$

If then the above series (31) is a solution of the differential equation (26), the expression (34) must be an identity; that is, must hold for every value of x .

If, however, it holds for every value of x , it does so also for $x=0$, and in this case, all the terms except the first vanish, and (34) becomes,

$$a_1 - a_0 = 0; \quad \text{or,} \quad a_1 = a_0. \quad (35)$$

To make (31) a solution of the differential equation $(a_1 - a_0)$ must therefore equal 0. This being the case, the term $(a_1 - a_0)$ can be dropped in (34), which then becomes,

$$(2a_2 - a_1)x + (3a_3 - a_2)x^2 + (4a_4 - a_3)x^3 + (5a_5 - a_4)x^4 + \dots = 0;$$

or,

$$x\{(2a_2 - a_1) + (3a_3 - a_2)x + (4a_4 - a_3)x^2 + \dots\} = 0.$$

Since this equation must hold for every value of x , the second factor of the equation must be zero, since the first factor, x , is not necessarily zero. This gives,

$$(2a_2 - a_1) + (3a_3 - a_2)x + (4a_4 - a_3)x^2 + \dots = 0.$$

As this equation holds for every value of x , it holds also for $x=0$. In this case, however, all terms except the first vanish, and,

$$2a_2 - a_1 = 0; \quad \dots \dots \dots (36)$$

hence,

$$a_2 = \frac{a_1}{2},$$

and from (35),

$$a_2 = \frac{a_0}{2}.$$

Continuing the same reasoning,

$$3a_3 - a_2 = 0, \quad 4a_4 - a_3 = 0, \text{ etc.}$$

Therefore, if an expression of successive powers of x , such as (34), is an identity, that is, holds for every value of x , then all the coefficients of all the powers of x must separately be zero.*

Hence,

$$\left. \begin{array}{l} a_1 - a_0 = 0; \quad \text{or} \quad a_1 = a_0; \\ 2a_2 - a_1 = 0; \quad \text{or} \quad a_2 = \frac{a_1}{2} = \frac{a_0}{2}; \\ 3a_3 - a_2 = 0; \quad \text{or} \quad a_3 = \frac{a_2}{3} = \frac{a_0}{3}; \\ 4a_4 - a_3 = 0; \quad \text{or} \quad a_4 = \frac{a_3}{4} = \frac{a_0}{4}; \\ \text{etc.,} \qquad \qquad \qquad \text{etc.} \end{array} \right\} \dots \dots \dots (37)$$

* The reader must realize the difference between an expression (34), as equation in x , and as substitution product of a function; that is, as an identity.

Regardless of the values of the coefficients, an expression (34) as equation gives a number of separate values of x , the roots of the equation, which make the left side of (34) equal zero, that is, solve the equation. If, however, the infinite series (31) is a solution of the differential equation (26), then the expression (34), which is the result of substituting (31) into (26), must be correct not only for a limited number of values of x , which are the roots of the equation, but for all values of x , that is, no matter what value is chosen for x , the left side of (34) must always give the same result, 0, that is, it must not be changed by a change of x , or in other words, it must not contain x , hence all the coefficients of the powers of x must be zero.

Therefore, if the coefficients of the series (31) are chosen by equation (37), this series satisfies the differential equation (26); that is,

$$y = a_0 \left\{ 1 + x + \frac{x^2}{2} + \frac{x^3}{\underline{3}} + \frac{x^4}{\underline{4}} + \dots \right\}. \quad (38)$$

is the solution of the differential equation,

$$\frac{dy}{dx} = y.$$

51. In the same manner, the differential equation (27),

$$\frac{dz}{dx} = az, \quad (39)$$

is solved by an infinite series,

$$z = a_0 + a_1x + a_2x^2 + a_3x^3 + \dots, \quad (40)$$

and the coefficients of this series determined by substituting (40) into (39), in the same manner as done above. This gives,

$$(a_1 - aa_0) + (2a_2 - aa_1)x + (3a_3 - aa_2)x^2 + (4a_4 - aa_3)x^3 + \dots = 0, \quad (41)$$

and, as this equation must be an identity, all its coefficients must be zero; that is,

$$\left. \begin{aligned} a_1 - aa_0 &= 0; \quad \text{or} \quad a_1 = aa_0; \\ 2a_2 - aa_1 &= 0; \quad \text{or} \quad a_2 = a_1 \frac{a}{2} = a_0 \frac{a^2}{2}; \\ 3a_3 - aa_2 &= 0; \quad \text{or} \quad a_3 = a_2 \frac{a}{3} = a_0 \frac{a^3}{\underline{3}}; \\ 4a_4 - aa_3 &= 0; \quad \text{or} \quad a_4 = a_3 \frac{a}{4} = a_0 \frac{a^4}{\underline{4}}; \\ \text{etc.,} & \qquad \qquad \qquad \text{etc.} \end{aligned} \right\}. \quad (42)$$

and the solution of differential equation (39) is,

$$z = a_0 \left\{ 1 + ax + \frac{a^2x^2}{2} + \frac{a^3x^3}{\underline{3}} + \frac{a^4x^4}{\underline{4}} + \dots \right\}. \quad (43)$$

52. These solutions, (38) and (43), of the differential equations (26) and (39), are not single solutions, but each contains an infinite number of solutions, as it contains an arbitrary

constant a_0 ; that is, a constant which may have any desired numerical value.

This can easily be seen, since, if z is a solution of the differential equation,

$$\frac{dz}{dx} = az,$$

then, any multiple, or fraction of z , bz , also is a solution of the differential equation;

$$\frac{d(bz)}{dx} = a(bz),$$

since the b cancels.

Such a constant, a_0 , which is not determined by the coefficients of the *mathematical* problem, but is left arbitrary, and requires for its determinations some further condition in addition to the differential equation, is called an *integration constant*. It usually is determined by some additional requirements of the *physical* problem, which the differential equation represents; that is, by a so-called *terminal condition*, as, for instance, by having the value of y given for some particular value of x , usually for $x=0$, or $x=\infty$.

The differential equation,

$$\frac{dy}{dx} = y; \quad (44)$$

thus, is solved by the function,

$$y = a_0 y_0. \quad (45)$$

where,

$$y_0 = 1 + x + \frac{x^2}{2} + \frac{x^3}{3} + \frac{x^4}{4} \quad (46)$$

and the differential equation,

$$\frac{dz}{dx} = az, \quad (47)$$

is solved by the function,

$$z = a_0 z_0. \quad (48)$$

where,

$$z_0 = 1 + ax + \frac{a^2 x^2}{2} + \frac{a^3 x^3}{3} + \frac{a^4 x^4}{4} \quad (49)$$

y_0 and z_0 thus are the simplest forms of the solutions y and z of the differential equations (26) and (39).

53. It is interesting now to determine the value of y^n . To raise the infinite series (46), which represents y_0 , to the n th power, would obviously be a very complicated operation.

However,

$$\frac{dy^n}{dx} = ny^{n-1} \frac{dy}{dx}, \quad . \quad . \quad . \quad . \quad . \quad . \quad (50)$$

and since from (44) $\frac{dy}{dx} = y, \quad . \quad . \quad . \quad . \quad . \quad . \quad (51)$

by substituting (51) into (50),

$$\frac{dy^n}{dx} = ny^n; \quad . \quad . \quad . \quad . \quad . \quad . \quad (52)$$

hence, the same equation as (47), but with y^n instead of z . Hence, if y is the solution of the differential equation,

$$\frac{dy}{dx} = y,$$

then $z = y^n$ is the solution of the differential equation (52),

$$\frac{dz}{dx} = nz.$$

However, the solution of this differential equation from (47), (48), and (49), is

$$z = a_0 z_0;$$

$$z_0 = 1 + nx + \frac{n^2 x^2}{2} + \frac{n^3 x^3}{|3|} + \dots;$$

that is, if

$$y_0 = 1 + x + \frac{x^2}{2} + \frac{x^3}{|3|} + \dots,$$

then,

$$z_0 = y_0^n = 1 + nx + \frac{n^2 x^2}{2} + \frac{n^3 x^3}{|3|} + \dots; \quad . \quad . \quad . \quad (53)$$

therefore the series y is raised to the n th power by multiplying the variable x by n .

Substituting now in equation (53) for n the value $\frac{1}{x}$ gives

$$y_0^{\frac{1}{x}} = 1 + 1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} + \dots; \quad (54)$$

that is, a constant numerical value. This numerical value equals 2.7182818. . . , and is usually represented by the symbol ϵ .

Therefore,

$$y_0^{\frac{1}{x}} = \epsilon;$$

hence,

$$y_0 = \epsilon^x = 1 + x + \frac{x^2}{2} + \frac{x^3}{3} + \frac{x^4}{4} + \dots \quad (55)$$

and

$$z_0 = y_0^n = (\epsilon^x)^n = \epsilon^{nx} = 1 + nx + \frac{n^2 x^2}{2} + \frac{n^3 x^3}{3} + \frac{n^4 x^4}{4} + \dots; \quad (56)$$

therefore, the infinite series, which integrates above differential equation, is an exponential function with the base

$$\epsilon = 2.7182818, \dots \quad (57)$$

The solution of the differential equation,

$$\frac{dy}{dx} = y, \quad (58)$$

thus is,

$$y = a_0 \epsilon^x, \quad (59)$$

and the solution of the differential equation,

$$\frac{dy}{dx} = ay, \quad (60)$$

is,

$$y = a_1 \epsilon^{ax}, \quad (61)$$

where a_0 is an integration constant.

The exponential function thus is one of the most common functions met in electrical engineering problems.

The above described method of solving a problem, by assuming a solution in a form containing a number of unknown coefficients, $a_0, a_1, a_2, \dots$, substituting the solution in the problem and thereby determining the coefficients, is called the *method of indeterminate coefficients*. It is one of the most convenient

and most frequently used methods of solving engineering problems.

EXAMPLE 1.

54. In a 4-pole 500-volt 50-kw. direct-current shunt motor, the resistance of the field circuit, inclusive of field rheostat, is 250 ohms. Each field pole contains 4000 turns, and produces at 500 volts impressed upon the field circuit, 8 megalines of magnetic flux per pole.

What is the equation of the field current, and how much time after closing the field switch is required for the field current to reach 90 per cent of its final value?

Let r be the resistance of the field circuit, L the inductance of the field circuit, and i the field current, then the voltage consumed in resistance is,

$$e_r = ri.$$

In general, in an electric circuit, the current produces a magnetic field; that is, lines of magnetic flux surrounding the conductor of the current; or, it is usually expressed, *interlinked with the current*. This magnetic field changes with a change of the current, and usually is proportional thereto. A change of the magnetic field surrounding a conductor, however, generates an e.m.f. in the conductor, and this e.m.f. is proportional to the rate of change of the magnetic field; hence, is proportional to the rate of change of the current, or to $\frac{di}{dt}$, with a proportionality factor L , which is called the *inductance* of the circuit. This counter-generated e.m.f. is in opposition to the current, $-L \frac{di}{dt}$, and thus consumes an e.m.f., $+L \frac{di}{dt}$, which is called the *e.m.f. consumed by self-inductance*, or *inductance e.m.f.*

Therefore, by the inductance, L , of the field circuit, a voltage is consumed which is proportional to the rate of change of the field current, thus,

$$e_x = L \frac{di}{dt}.$$

Since the supply voltage, and thus the total voltage consumed in the field circuit, is $e=500$ volts,

$$e=ri+L\frac{di}{dt}; \quad . \quad . \quad . \quad . \quad . \quad . \quad (62)$$

or, rearranged,

$$\frac{di}{dt}=\frac{e-ri}{L}.$$

Substituting herein,

$$u=e-ri; \quad . \quad . \quad . \quad . \quad . \quad . \quad (63)$$

hence,

$$\frac{du}{dt}=-r\frac{di}{dt},$$

gives,

$$\frac{du}{dt}=-\frac{r}{L}u. \quad . \quad . \quad . \quad . \quad . \quad . \quad (64)$$

This is the same differential equation as (39), with $a=-\frac{r}{L}$, and therefore is integrated by the function,

$$u=a_0e^{-\frac{r}{L}t};$$

therefore, resubstituting from (63),

$$e-ri=a_0e^{-\frac{r}{L}t},$$

and

$$i=\frac{e}{r}-\frac{a_0}{r}e^{-\frac{r}{L}t}. \quad . \quad . \quad . \quad . \quad . \quad . \quad (65)$$

This solution (65), still contains the unknown quantity a_0 ; or, the integration constant, and this is determined by knowing the current i for some particular value of the time t .

Before closing the field switch and thereby impressing the voltage on the field, the field current obviously is zero. In the moment of closing the field switch, the current thus is still zero; that is,

$$i=0 \text{ for } t=0. \quad . \quad . \quad . \quad . \quad . \quad . \quad (66)$$

Substituting these values in (65) gives,

$$0 = \frac{e}{r} - \frac{a_0}{r}; \quad \text{or} \quad a_0 = +e,$$

hence,

$$i = \frac{e}{r} \left(1 - \varepsilon^{-\frac{r}{L}t} \right) \quad . \quad . \quad . \quad . \quad . \quad (67)$$

is the final solution of the differential equation (62); that is, it is the value of the field current, i , as function of the time, t , after closing the field switch.

After infinite time, $t = \infty$, the current i assumes the final value i_0 , which is given by substituting $t = \infty$ into equation (67), thus,

$$i_0 = \frac{e}{r} = \frac{500}{250} = 2 \text{ amperes}; \quad . \quad . \quad . \quad . \quad . \quad (68)$$

hence, by substituting (68) into (67), this equation can also be written,

$$\begin{aligned} i &= i_0 \left(1 - \varepsilon^{-\frac{r}{L}t} \right) \\ &= 2 \left(1 - \varepsilon^{-\frac{r}{L}t} \right), \quad . \quad . \quad . \quad . \quad . \quad (69) \end{aligned}$$

where $i_0 = 2$ is the final value assumed by the field current.

The time t_1 , after which the field current i has reached 90 per cent of its final value i_0 , is given by substituting $i = 0.9i_0$ into (69), thus,

$$0.9i_0 = i_0 \left(1 - \varepsilon^{-\frac{r}{L}t_1} \right),$$

and

$$\varepsilon^{-\frac{r}{L}t_1} = 0.1.$$

Taking the logarithm of both sides,

$$-\frac{r}{L} t_1 \log \varepsilon = -1;$$

and

$$t_1 = \frac{L}{r \log \varepsilon} \quad . \quad . \quad . \quad . \quad . \quad (70)$$

55. The inductance L is calculated from the data given in the problem. Inductance is measured by the number of interlinkages of the electric circuit, with the magnetic flux produced by one absolute unit of current in the circuit; that is, it equals the product of magnetic flux and number of turns divided by the absolute current.

A current of $i_0 = 2$ amperes represents 0.2 absolute units, since the absolute unit of current is 10 amperes. The number of field turns per pole is 4000; hence, the total number of turns $n = 4 \times 4000 = 16,000$. The magnetic flux at full excitation, or $i_0 = 0.2$ absolute units of current, is given as $\Phi = 8 \times 10^6$ lines of magnetic force. The inductance of the field thus is:

$$L = \frac{n\Phi}{i_0} = \frac{16000 \times 8 \times 10^6}{0.2} = 640 \times 10^9 \text{ absolute units} = 640h,$$

the practical unit of inductance, or the henry (h) being 10^9 absolute units.

Substituting $L = 640$ $r = 250$ and $c = 500$, into equation (67) and (70) gives

$$i = 2(1 - e^{-0.39t}),$$

and

$$t_1 = \frac{640}{250 \times 0.4343} = 5.88 \text{ sec.} \quad (71)$$

Therefore it takes about 6 sec. before the motor field has reached 90 per cent of its final value.

The reader is advised to calculate and plot the numerical values of i from equation (71), for

$t = 0, 0.1, 0.2, 0.4, 0.6, 0.8, 1.0, 1.5, 2.0, 3, 4, 5, 6, 8, 10$ sec.

This calculation is best made in the form of a table, thus;

$$e^{-0.39t} = N - 0.39t \log e,$$

and,

$$\log e = 0.4343;$$

hence,

$$0.39t \log e = 0.1694t;$$

and,

$$e^{-0.39t} = N - 0.1694t.$$

The values of $\varepsilon^{-0.39t}$ can also be taken directly from the tables of the exponential function, at the end of the book.

t	$0.1694t$	$-0.1694t$	$\varepsilon^{-0.39t}$ $= N^{-0.1694t}$	$1 - \varepsilon^{-0.39t}$	$i =$ $2(1 - \varepsilon^{-0.39t})$
0.0	0	0	1	0	0
0.1	0.0170	0.9830-1	0.962	0.038	0.076
0.2	0.0339	0.9661-1	0.925	0.075	0.150
0.4	0.0678	0.9322-1	0.855	0.145	0.290
0.6	0.1016	0.8984-1	0.791	0.209	0.418
0.8	0.1355	0.8645-1	0.732	0.268	0.536
etc.					
....					

EXAMPLE 2.

56. A condenser of 20 mf. capacity, is charged to a potential of $e_0 = 10,000$ volts, and then discharges through a resistance of 2 megohms. What is the equation of the discharge current,

and after how long a time has the voltage at the condenser dropped to 0.1 its initial value?

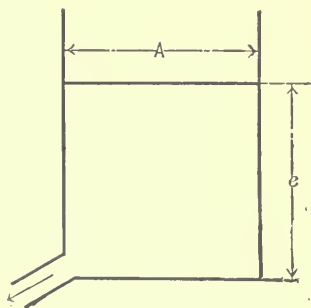


FIG. 27. Water Reservoir.

A condenser acts as a reservoir of electric energy, similar to a tank as water reservoir. If in a water tank, Fig. 27, A is the sectional area of the tank, e , the height of water, or water pressure, and water flows out of the tank, then the height e decreases by the flow of water; that is the tank empties, and

the current of water, i , is proportional to the change of the water level or height of water, $\frac{de}{dt}$, and to the area A of the tank; that is, it is,

$$i = -A \frac{de}{dt}. \quad \dots \dots \dots (72)$$

The minus sign stands on the right-hand side, as for positive i ; that is, out-flow, the height of the water decreases; that is, de is negative.

In an electric reservoir, the electric pressure or voltage e corresponds to the water pressure or height of the water, and to the storage capacity or sectional area A of the water tank corresponds the electric storage capacity of the condenser, called capacity C . The current or flow out of an electric condenser, thus is,

$$i = -C \frac{de}{dt}. \quad \dots \dots \dots (73)$$

The capacity of condenser is,

$$C = 20 \text{ mf} = 20 \times 10^{-6} \text{ farads.}$$

The resistance of the discharge path is,

$$r = 2 \times 10^6 \text{ ohms;}$$

hence, the current taken by the resistance, r , is

$$i = \frac{e}{r},$$

and thus

$$-C \frac{de}{dt} = \frac{e}{r};$$

and

$$\frac{de}{dt} = -\frac{1}{Cr} e.$$

Therefore, from (60) (61),

$$e = a_0 \varepsilon^{-\frac{t}{Cr}},$$

and for $t=0$, $e=e_0=10,000$ volts; hence

$$10,000 = a_0,$$

and

$$\begin{aligned} e &= e_0 \varepsilon^{-\frac{t}{Cr}} \\ &= 10,000 \varepsilon^{-0.025t} \text{ volts; } \dots \dots \dots (74) \end{aligned}$$

0.1 of the initial value:

$$e = 0.1 e_0,$$

is reached at:

$$t_1 = \frac{Cr}{\log \varepsilon} = 92 \text{ sec. } \dots \dots \dots (75)$$

Therefore

$$\begin{aligned}
 a_2 &= \frac{aa_0}{2}; & a_3 &= \frac{aa_1}{3}; \\
 a_4 &= \frac{aa_2}{3 \times 4} = \frac{a_0 a^2}{4}; & a_5 &= \frac{aa_3}{4 \times 5} = \frac{a_1 a^2}{5}; \\
 a_6 &= \frac{aa_4}{5 \times 6} = \frac{a_0 a^3}{6}; & a_7 &= \frac{aa_5}{6 \times 7} = \frac{a_1 a^3}{7}; \\
 a_8 &= \frac{aa_6}{7 \times 8} = \frac{a_0 a^4}{8}; & a_9 &= \frac{aa_7}{8 \times 9} = \frac{a_1 a^4}{9}; \\
 &\text{etc.,} & &\text{etc.}
 \end{aligned}$$

Substituting these values in (77),

$$\begin{aligned}
 y = a_0 \left\{ 1 + \frac{ax^2}{2} + \frac{a^2 x^4}{4} + \frac{a^3 x^6}{6} + \dots \right\} \\
 + a_1 \left\{ x + \frac{ax^3}{3} + \frac{a^2 x^5}{5} + \frac{a^3 x^7}{7} + \dots \right\}. \quad (80)
 \end{aligned}$$

In this case, two coefficients a_0 and a_1 thus remain indeterminate, as was to be expected, as a differential equation of second order must have two integration constants in its most general form of solution.

Substituting into this equation,

$$b^2 = a;$$

that is,

$$b = \sqrt{a}. \quad (81)$$

$$\frac{d^2 y}{dx^2} = +b^2 y. \quad (82)$$

and

$$\begin{aligned}
 y = a_0 \left\{ 1 + \frac{b^2 x^2}{2} + \frac{b^4 x^4}{4} + \frac{b^6 x^6}{6} + \dots \right\} \\
 + \frac{a_1}{b} \left\{ bx + \frac{b^3 x^3}{3} + \frac{b^5 x^5}{5} + \frac{b^7 x^7}{7} + \dots \right\}. \quad (83)
 \end{aligned}$$

In this case, instead of the integration constants a_0 and a_1 , the two new integration constants A and B can be introduced by the equations

$$a_0 = A + B \quad \text{and} \quad \frac{a_1}{b} = A - B;$$

hence,

$$A = \frac{a_0 + \frac{a_1}{b}}{2} \quad \text{and} \quad B = \frac{a_0 - \frac{a_1}{b}}{2},$$

and, substituting these into equation (83), gives,

$$y = A \left\{ 1 + bx + \frac{b^2 x^2}{|2|} + \frac{b^3 x^3}{|3|} + \frac{b^4 x^4}{|4|} + \dots \right\} \\ + B \left\{ 1 - bx + \frac{b^2 x^2}{|2|} - \frac{b^3 x^3}{|3|} + \frac{b^4 x^4}{|4|} - + \dots \right\} \quad (84)$$

The first series, however, from (56), for $n=b$ is ϵ^{+bx} , and the second series from (56), for $n=-b$ is ϵ^{-bx} .

Therefore, the infinite series (83) is,

$$y = A\epsilon^{+bx} + B\epsilon^{-bx}; \quad (85)$$

that is, it is the sum of two exponential functions, the one with a positive, the other with a negative exponent.

Hence, the differential equation,

$$\frac{d^2 y}{dx^2} = ay, \quad (76)$$

is integrated by the function,

$$y = A\epsilon^{+bx} + B\epsilon^{-bx}, \quad (86)$$

where,

$$b = \sqrt{a}. \quad (87)$$

However, if a is a negative quantity, $b = \sqrt{a}$ is imaginary, and can be represented by

$$b = jc, \quad (88)$$

where

$$c^2 = -a. \quad (89)$$

In this case, equation (86) assumes the form,

$$y = A\epsilon^{+jcx} + B\epsilon^{-jcx}; \quad (90)$$

that is, if in the differential equation (76) a is a positive quantity, $= +b^2$, this differential equation is integrated by the sum of the two exponential functions (86); if, however, a is a negative quantity, $= -c^2$, the solution (86) appears in the form of exponential functions with imaginary exponents (90).

58. In the latter case, a form of the solution of differential equation (76) can be derived which does not contain the imaginary appearance, by turning back to equation (80), and substituting therein $a = -c^2$, which gives,

$$\frac{d^2y}{dx^2} = -c^2y \quad . \quad . \quad . \quad . \quad . \quad . \quad (91)$$

$$y = a_0 \left\{ 1 - \frac{c^2x^2}{2} + \frac{c^4x^4}{|4|} - \frac{c^6x^6}{|6|} + \dots \right\} - \frac{a_1}{c} \left\{ cx - \frac{c^3x^3}{|3|} + \frac{c^5x^5}{|5|} - + \dots \right\};$$

or, writing $A = a_0$ and $B = -\frac{a_1}{c}$,

$$y = A \left\{ 1 - \frac{c^2x^2}{2} + \frac{c^4x^4}{|4|} - \frac{c^6x^6}{|6|} + \dots \right\} + B \left\{ cx - \frac{c^3x^3}{|3|} + \frac{c^5x^5}{|5|} - + \dots \right\}. \quad (92)$$

The solution then is given by the sum of two infinite series, thus,

$$\left. \begin{aligned} u(cx) &= 1 - \frac{c^2x^2}{2} + \frac{c^4x^4}{|4|} - \frac{c^6x^6}{|6|} + \dots \\ v(cx) &= cx - \frac{c^3x^3}{|3|} + \frac{c^5x^5}{|5|} - + \dots \end{aligned} \right\} \quad . \quad . \quad (93)$$

as

$$y = Au(cx) + Bv(cx). \quad . \quad . \quad . \quad . \quad . \quad . \quad (94)$$

In the u -series, a change of the sign of x does not change the value of u ,

$$u(-cx) = u(+cx). \quad . \quad . \quad . \quad . \quad . \quad . \quad (95)$$

Such a function is called an even function.

In the v -series, a change of the sign of x reverses the sign of v , as seen from (93):

$$v(-cx) = -v(cx). \quad . \quad . \quad . \quad . \quad . \quad (96)$$

Such a function is called an odd function.

It can be shown that

$$u(cx) = \cos cx \quad \text{and} \quad v(cx) = \sin cx; \quad . \quad . \quad . \quad (97)$$

hence,

$$y = A \cos cx + B \sin cx, \quad . \quad . \quad . \quad . \quad . \quad (98)$$

where A and B are the integration constants, which are to be determined by the terminal conditions of the physical problem.

Therefore, the solution of the differential equation

$$\frac{d^2y}{dx^2} = ay, \quad . \quad . \quad . \quad . \quad . \quad (99)$$

has two different forms, an exponential and a trigonometric.

If a is positive,

$$\frac{d^2y}{dx^2} = +b^2y, \quad . \quad . \quad . \quad . \quad . \quad (100)$$

it is:

$$y = A e^{+bx} + B e^{-bx}, \quad . \quad . \quad . \quad . \quad . \quad (101)$$

If a is negative,

$$\frac{d^2y}{dx^2} = -c^2y, \quad . \quad . \quad . \quad . \quad . \quad (102)$$

it is:

$$y = A \cos cx + B \sin cx. \quad . \quad . \quad . \quad . \quad . \quad (103)$$

In the latter case, the solution (101) would appear as exponential function with imaginary exponents;

$$y = A e^{+icx} + B e^{-icx} \quad . \quad . \quad . \quad . \quad . \quad (104)$$

As (104) obviously must be the same function as (103), it follows that exponential functions with imaginary exponents must be expressible by trigonometric functions.

59. The exponential functions and the trigonometric functions, according to the preceding discussion, are expressed by the infinite series,

$$\left. \begin{aligned} e^x &= 1 + x + \frac{x^2}{2} + \frac{x^3}{3} + \frac{x^4}{4} + \frac{x^5}{5} + \dots \\ \cos x &= 1 - \frac{x^2}{2} + \frac{x^4}{4} - \frac{x^6}{6} + \dots \\ \sin x &= x - \frac{x^3}{3} + \frac{x^5}{5} - \frac{x^7}{7} + \dots \end{aligned} \right\} \dots \dots (105)$$

Therefore, substituting ju for x ,

$$\begin{aligned} e^{ju} &= 1 + ju - \frac{u^2}{2} - j\frac{u^3}{3} + \frac{u^4}{4} + j\frac{u^5}{5} - \frac{u^6}{6} - j\frac{u^7}{7} + \dots \\ &= \left(1 - \frac{u^2}{2} + \frac{u^4}{4} - \frac{u^6}{6} + \dots\right) + j\left(u - \frac{u^3}{3} + \frac{u^5}{5} - \frac{u^7}{7} + \dots\right) \end{aligned}$$

However, the first part of this series is $\cos u$, the latter part $\sin u$, by (105); that is,

$$e^{ju} = \cos u + j \sin u. \quad (106)$$

Substituting $-u$ for $+u$ gives,

$$e^{-ju} = \cos u - j \sin u \quad (107)$$

Combining (106) and (107) gives,

$$\left. \begin{aligned} \cos u &= \frac{e^{+ju} + e^{-ju}}{2} \\ \sin u &= \frac{e^{+ju} - e^{-ju}}{2j} \end{aligned} \right\} \dots \dots (108)$$

Substituting in (106) to (108), jr for u , gives,

$$\left. \begin{aligned} e^{-jr} &= \cos jr + j \sin jr \\ e^{+jr} &= \cos jr - j \sin jr \end{aligned} \right\} \dots \dots (109)$$

Adding and subtracting gives respectively,

$$\left. \begin{aligned} \cos jv &= \frac{\epsilon^{+v} + \epsilon^{-v}}{2}, \\ \sin jv &= \frac{\epsilon^{-v} - \epsilon^{+v}}{2j}. \end{aligned} \right\} \dots \dots (110)$$

By these equations, (106) to (110), exponential functions with imaginary exponents can be transformed into trigonometric functions with real angles, and exponential functions with real exponents into trigonometric functions with imaginary angles, and inversely.

Mathematically, the trigonometric functions thus do not constitute a separate class of functions, but may be considered as exponential functions with imaginary angles, and it can be said broadly that the solution of the above differential equations is given by the exponential function, but that in this function the exponent may be real, or may be imaginary, and in the latter case, the expression is put into real form by introducing the trigonometric functions.

EXAMPLE 1.

60. A condenser (as an underground high-potential cable) of 20 mf. capacity, and of a voltage of $e_0 = 10,000$, discharges through an inductance of 50 mh. and of negligible resistance, What is the equation of the discharge current?

The current consumed by a condenser of capacity C and potential difference e is proportional to the rate of change of the potential difference, and to the capacity; hence, it is $C \frac{de}{dt}$, and the current from the condenser; or its discharge current, is

$$i = -C \frac{de}{dt} \dots \dots (111)$$

The voltage consumed by an inductance L is proportional to the rate of change of the current in the inductance, and to the inductance; hence,

$$e = L \frac{di}{dt} \dots \dots (112)$$

Differentiating (112) gives,

$$\frac{de}{dt} = L \frac{d^2 i}{dt^2},$$

and substituting this into (111) gives,

$$i = -CL \frac{d^2 i}{dt^2}; \quad \text{or,} \quad \frac{d^2 i}{dt^2} = -\frac{1}{CL} i, \quad . \quad . \quad . \quad (113)$$

as the differential equation of the problem.

This equation (113) is the same as (102), for $c^2 = \frac{1}{CL}$, thus is solved by the expression,

$$i = A \cos \frac{t}{\sqrt{LC}} + B \sin \frac{t}{\sqrt{LC}}, \quad . \quad . \quad . \quad (114)$$

and the potential difference at the condenser or at the inductance is, by substituting (114) into (112),

$$e = \sqrt{\frac{L}{C}} \left\{ B \cos \frac{t}{\sqrt{LC}} - A \sin \frac{t}{\sqrt{LC}} \right\}. \quad . \quad . \quad (115)$$

These equations (114) and (115) still contain two unknown constants, A and B , which have to be determined by the terminal conditions, that is, by the known conditions of current and voltage at some particular time.

At the moment of starting the discharge; or, at the time $t=0$, the current is zero, and the voltage is that to which the condenser is charged, that is, $i=0$, and $e=e_0$.

Substituting these values in equations (114) and (115) gives,

$$0 = A \quad \text{and} \quad e_0 = \sqrt{\frac{L}{C}} B;$$

hence

$$B = e_0 \sqrt{\frac{C}{L}},$$

and, substituting for A and B the values in (114) and (115), gives

$$\left. \begin{aligned} i &= e_0 \sqrt{\frac{C}{L}} \sin \frac{t}{\sqrt{CL}}, \\ e &= e_0 \cos \frac{t}{\sqrt{CL}} \end{aligned} \right\} . \quad . \quad . \quad . \quad (116)$$

and

Substituting the numerical values, $e_0 = 10,000$ volts, $C = 20$ mf. $= 20 \times 10^{-6}$ farads, $L = 50$ mh. $= 0.05$ h. gives,

$$\sqrt{\frac{C}{L}} = 0.02 \quad \text{and} \quad \sqrt{CL} = 10^{-3};$$

hence,

$$i = 200 \sin 1000 t \quad \text{and} \quad e = 10,000 \cos 1000 t.$$

61. The discharge thus is alternating. In reality, due to the unavoidable resistance in the discharge path, the alternations gradually die out, that is, the discharge is oscillating.

The time of one complete period is given by,

$$1000t_0 = 2\pi; \quad \text{or,} \quad t_0 = \frac{2\pi}{1000}.$$

Hence the frequency,

$$f = \frac{1}{t_0} = \frac{1000}{2\pi} = 159 \text{ cycles per second.}$$

As the circuit in addition to the inductance necessarily contains resistance r , besides the voltage consumed by the inductance by equation (112), voltage is consumed by the resistance, thus

$$e_r = ri, \quad . \quad . \quad . \quad . \quad . \quad . \quad (117)$$

and the total voltage consumed by resistance r and inductance L , thus is

$$e = ri + L \frac{di}{dt}. \quad . \quad . \quad . \quad . \quad . \quad . \quad (118)$$

Differentiating (118) gives,

$$\frac{de}{dt} = r \frac{di}{dt} + L \frac{d^2i}{dt^2}, \quad . \quad . \quad . \quad . \quad . \quad . \quad (119)$$

and, substituting this into equation (111), gives,

$$i + Cr \frac{di}{dt} + CL \frac{d^2i}{dt^2} = 0, \quad . \quad . \quad . \quad . \quad . \quad . \quad (120)$$

as the differential equation of the problem.

This differential equation is of the more general form, (30),

62. The more general differential equation (30).

$$\frac{d^2y}{dx^2} + 2c \frac{dy}{dx} + ay + b = 0, \quad . \quad . \quad . \quad . \quad . \quad . \quad (121)$$

can, by substituting,

$$y + \frac{b}{a} = z, \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (122)$$

which gives

$$\frac{dy}{dx} = \frac{dz}{dx},$$

be transformed into the somewhat simpler form,

$$\frac{d^2z}{dx^2} + 2c \frac{dz}{dx} + az = 0. \quad . \quad . \quad . \quad . \quad . \quad . \quad (123)$$

It may also be solved by the method of indeterminate coefficients, by substituting for z an infinite series of powers of x , and determining thereby the coefficients of the series.

As, however, the simpler forms of this equation were solved by exponential functions, the applicability of the exponential functions to this equation (123) may be directly tried, by the method of indeterminate coefficients. That is, assume as solution an exponential function,

$$z = A\epsilon^{bx}, \quad . \quad . \quad . \quad . \quad . \quad . \quad (124)$$

where A and b are unknown constants. Substituting (124) into (123), if such values of A and b can be found, which make the substitution product an identity, (124) is a solution of the differential equation (123).

From (124) it follows that,

$$\frac{dz}{dx} = bA\epsilon^{bx}; \quad \text{and} \quad \frac{d^2z}{dx^2} = b^2A\epsilon^{bx}, \quad . \quad . \quad . \quad (125)$$

and substituting (124) and (125) into (123), gives,

$$A\epsilon^{bx}(b^2 + 2cb + a) = 0. \quad . \quad . \quad . \quad . \quad (126)$$

As seen, this equation is satisfied for every value of x , that is, it is an identity, if the parenthesis is zero, thus,

$$b^2 + 2cb + a = 0, \quad . \quad . \quad . \quad . \quad (127)$$

and the value of b , calculated by the quadratic equation (127), thus makes (124) a solution of (123), and leaves A still undetermined, as integration constant.

From (127),

$$b = -c \pm \sqrt{c^2 - a};$$

or, substituting,

$$\sqrt{c^2 - a} = p, \quad . \quad . \quad . \quad . \quad . \quad (128)$$

into (128), the equation becomes,

$$b = -c \pm p. \quad . \quad . \quad . \quad . \quad . \quad (129)$$

Hence, two values of b exist,

$$b_1 = -c + p \quad \text{and} \quad b_2 = -c - p,$$

and, therefore, the differential equation,

$$\frac{d^2z}{dx^2} + 2c \frac{dz}{dx} + az = 0, \quad . \quad . \quad . \quad . \quad . \quad (130)$$

is solved by Ae^{b_1x} ; or, by Ae^{b_2x} , or, by any combination of these two solutions. The most general solution is,

$$z = A_1 e^{b_1x} + A_2 e^{b_2x};$$

that is,

$$\left. \begin{aligned} y &= A_1 e^{(-c+p)x} + A_2 e^{(-c-p)x} - \frac{b}{a} \\ &= e^{-cx} \{ A_1 e^{+px} + A_2 e^{-px} \} - \frac{b}{a} \end{aligned} \right\} . \quad . \quad . \quad (131)$$

As roots of a quadratic equation, b_1 and b_2 may both be real quantities, or may be complex imaginary, and in the latter case, the solution (131) appears in imaginary form, and has to be reduced or modified for use, so as to eliminate the imaginary appearance, by the relations (106) and (107).

EXAMPLE 2.

63. Assume, in the example in paragraph 60, the discharge circuit of the condenser of $C=20$ mf. capacity, to contain, besides the inductance, $L=0.05$ h, the resistance, $r=125$ ohms.

The general equation of the problem, (120), dividing by $C L$, becomes,

$$\frac{d^2i}{dt^2} + \frac{r}{L} \frac{di}{dt} + \frac{i}{CL} = 0. \quad . \quad . \quad . \quad . \quad . \quad (132)$$

This is the equation (123), for:

$$\left. \begin{aligned} x=t, \quad 2c=\frac{r}{L}=2500; \\ z=i, \quad a=\frac{1}{CL}=10^6 \end{aligned} \right\} \quad . \quad . \quad . \quad . \quad (133)$$

If $p=\sqrt{c^2-a}$, then

$$\left. \begin{aligned} p &= \sqrt{\left(\frac{r}{2L}\right)^2 - \frac{1}{CL}} \\ &= \frac{1}{2L} \sqrt{r^2 - 4\frac{L}{C}}, \end{aligned} \right\} \quad . \quad . \quad . \quad . \quad (134)$$

and, writing

$$\left. \begin{aligned} \sqrt{r^2 - \frac{4L}{C}} &= s, \\ p &= \frac{s}{2L}, \end{aligned} \right\}, \quad . \quad . \quad . \quad . \quad (135)$$

and since

$$\left. \begin{aligned} \frac{1}{2L} &= 10 \quad \text{and} \quad \frac{L}{C} = 2500, \\ s &= 75 \quad \text{and} \quad p = 750. \end{aligned} \right\} \quad . \quad . \quad . \quad . \quad (136)$$

The equation of the current from (131) then is,

$$\left. \begin{aligned} i &= A_1 \varepsilon^{-\frac{r-a}{2L}t} + A_2 \varepsilon^{-\frac{r+a}{2L}t} \\ &= \varepsilon^{-\frac{r}{2L}t} \left\{ A_1 \varepsilon^{+\frac{a}{2L}t} + A_2 \varepsilon^{-\frac{a}{2L}t} \right\}. \end{aligned} \right\} \quad . \quad . \quad . \quad (137)$$

This equation still contains two unknown quantities, the integration constants A_1 and A_2 , which are determined by the terminal condition: The values of current and of voltage at the beginning of the discharge, or $t=0$.

This requires the determination of the equation of the voltage at the condenser terminals. This obviously is the voltage consumed by resistance and inductance, and is expressed by equation (118),

$$e = ri + L \frac{di}{dt}; \quad . \quad . \quad . \quad . \quad (118)$$

hence, substituting herein the value of i and $\frac{di}{dt}$, from equation (137), gives

$$\begin{aligned} e &= r \left\{ A_1 \varepsilon^{-\frac{r-s}{2L}t} + A_2 \varepsilon^{-\frac{r+s}{2L}t} \right\} + L \left\{ -\frac{r-s}{2L} A_1 \varepsilon^{-\frac{r-s}{2L}t} - \frac{r+s}{2L} A_2 \varepsilon^{-\frac{r+s}{2L}t} \right\} \\ &= \frac{r+s}{2} A_1 \varepsilon^{-\frac{r-s}{2L}t} + \frac{r-s}{2} A_2 \varepsilon^{-\frac{r+s}{2L}t} \\ &= \varepsilon^{-\frac{r}{2L}t} \left\{ \frac{r+s}{2} A_1 \varepsilon^{+\frac{s}{2L}t} + \frac{r-s}{2} A_2 \varepsilon^{-\frac{s}{2L}t} \right\}, \quad \dots \quad (138) \end{aligned}$$

and, substituting the numerical values (133) and (136) into equations (137) and (138), gives

$$\begin{aligned} i &= A_1 \varepsilon^{-500t} + A_2 \varepsilon^{-2000t} \\ \text{and,} \quad e &= 100 A_1 \varepsilon^{-500t} + 25 A_2 \varepsilon^{-2000t} \end{aligned} \quad \left. \vphantom{\begin{aligned} i &= A_1 \varepsilon^{-500t} + A_2 \varepsilon^{-2000t} \\ e &= 100 A_1 \varepsilon^{-500t} + 25 A_2 \varepsilon^{-2000t} \end{aligned}} \right\} \dots \quad (139)$$

At the moment of the beginning of the discharge, $t=0$, the current is zero and the voltage is 10,000; that is,

$$t=0; \quad i=0; \quad e=10,000 \quad \dots \quad (140)$$

Substituting (140) into (139) gives,

$$0 = A_1 + A_2, \quad 10,000 = 100A_1 + 25A_2;$$

hence,

$$A_2 = -A_1; \quad A_1 = 133.3; \quad A_2 = -133.3. \quad \dots \quad (141)$$

Therefore, the current and voltage are,

$$\begin{aligned} i &= 133.3 \left\{ \varepsilon^{-500t} - \varepsilon^{-2000t} \right\}, \\ e &= 13,333 \varepsilon^{-500t} - 3333 \varepsilon^{-2000t} \end{aligned} \quad \left. \vphantom{\begin{aligned} i &= 133.3 \{ \varepsilon^{-500t} - \varepsilon^{-2000t} \} \\ e &= 13,333 \varepsilon^{-500t} - 3333 \varepsilon^{-2000t} \end{aligned}} \right\} \dots \quad (142)$$

The reader is advised to calculate and plot the numerical values of i and e , and of their two components, for,

$$t = 0, 0.2, 0.4, 0.6, 1, 1.2, 1.5, 2, 2.5, 3, 4, 5, 6 \times 10^{-3} \text{ sec.}$$

64. Assuming, however, that the resistance of the discharge circuit is only $r=80$ ohms (instead of 125 ohms, as assumed above):

$r^2 - \frac{4L}{C}$ in equation (134) then becomes -3600 , and therefore:

$$s = \sqrt{-3600} = 60\sqrt{-1} = 60j,$$

and

$$p = \frac{s}{2L} = 600j.$$

The equation of the current (137) thus appears in imaginary form,

$$i = e^{-800t} \{ A_1 e^{+600jt} + A_2 e^{-600jt} \}. \quad (143)$$

The same is also true of the equation of voltage.

As it is obvious, however, physically, that a real current must be coexistent with a real e.m.f., it follows that this imaginary form of the expression of current and voltage is only apparent, and that in reality, by substituting for the exponential functions with imaginary exponents their trigonometric expressions, the imaginary terms must eliminate, and the equation (143) appear in real form.

According to equations (106) and (107),

$$\left. \begin{aligned} e^{+600jt} &= \cos 600t + j \sin 600t; \\ e^{-600jt} &= \cos 600t - j \sin 600t. \end{aligned} \right\} \quad (144)$$

Substituting (144) into (143) gives,

$$i = e^{-800t} \{ B_1 \cos 600t + B_2 \sin 600t \}, \quad (145)$$

where B_1 and B_2 are combinations of the previous integration constants A_1 and A_2 thus,

$$B_1 = A_1 + A_2, \quad \text{and} \quad B_2 = j(A_1 - A_2). \quad (146)$$

By substituting the numerical values, the condenser e.m.f., given by equation (138), then becomes,

$$\begin{aligned} e &= e^{-800t} \{ (40 + 30j)A_1(\cos 600t + j \sin 600t) \\ &\quad + (40 - 30j)A_2(\cos 600t - j \sin 600t) \} \\ &= e^{-800t} \{ (40B_1 + 30B_2)\cos 600t + (40B_2 - 30B_1)\sin 600t \}. \end{aligned} \quad (147)$$

Since for $t=0$, $i=0$ and $e=10,000$ volts (140), substituting into (145) and (147),

$$0 = B_1 \text{ and } 10,000 = 40 B_1 + 30 B_2.$$

Therefore, $B_1=0$ and $B_2=333$ and, by (145) and (147),

$$\left. \begin{aligned} i &= 333\epsilon^{-800t} \sin 600 t; \\ e &= 10,000\epsilon^{-800t} (\cos 600 t + 1.33 \sin 600 t). \end{aligned} \right\} \dots (148)$$

As seen, in this case the current i is larger, and current and e.m.f. are the product of an exponential term (gradually decreasing value) and a trigonometric term (alternating value); that is, they consist of successive alternations of gradually decreasing amplitude. Such functions are called *oscillating functions*. Practically all disturbances in electric circuits consist of such oscillating currents and voltages.

$600t=2\pi$ gives, as the time of one complete period,

$$T = \frac{2\pi}{600} = 0.0105 \text{ sec.};$$

and the frequency is

$$f = \frac{1}{T} = 95.3 \text{ cycles per sec.}$$

In this particular case, as the resistance is relatively high, the oscillations die out rather rapidly.

The reader is advised to calculate and plot the numerical values of i and e , and of their exponential terms, for every 30 degrees, that is, for $t=0$, $\frac{T}{12}$, $2\frac{T}{12}$, $3\frac{T}{12}$, etc., for the first two periods, and also to derive the equations, and calculate and plot the numerical values, for the same capacity, $C=20$ mf., and same inductance, $L=0.05h$, but for the much lower resistance, $r=20$ ohms.

65. Tables of ϵ^{+x} and ϵ^{-x} , for 5 decimals, and tables of $\log \epsilon^{+x}$ and $\log \epsilon^{-x}$, for 6 decimals, are given at the end of the book, and also a table of ϵ^{-x} for 3 decimals. For most engineering purposes the latter is sufficient; where a higher accuracy is required, the 5 decimal table may be used, and for

highest accuracy interpolation by the logarithmic table may be employed. For instance,

$$\epsilon^{-13.6847} = ?$$

From the logarithmic table,

$$\log \epsilon^{-10} = 5.657055,$$

$$\log \epsilon^{-3} = 8.697117,$$

$$\log \epsilon^{-0.6} = 9.739423,$$

$$\log \epsilon^{-0.08} = 9.965256,$$

$$\log \epsilon^{-0.0047} = 9.997959,$$

added

$$\left\{ \begin{array}{l} \text{interpolated,} \\ \text{between } \log \epsilon^{-0.004} = 9.998263, \\ \text{and } \log \epsilon^{-0.005} = 9.997829, \end{array} \right.$$

$$\log \epsilon^{-13.6847} = 4.056810 = 0.056810 - 6.$$

From common logarithmic tables,

$$\epsilon^{-13.6847} = 1.13976 \times 10^{-6}.$$

NOTE. In mathematics, for the base of the natural logarithms, 2.718282 . . . , is usually chosen the symbol e . Since, however, in engineering the symbol e is universally used to represent *voltage*, for the base of natural logarithms has been chosen the symbol ϵ , as the Greek letter corresponding to e , and ϵ is generally used in electrical engineering calculations in this meaning.

CHAPTER III.

TRIGONOMETRIC SERIES.

A. TRIGONOMETRIC FUNCTIONS.

66. For the engineer, and especially the electrical engineer, a perfect familiarity with the trigonometric functions and trigonometric formulas is almost as essential as familiarity with the multiplication table. To use trigonometric methods efficiently, it is not sufficient to understand trigonometric formulas enough to be able to look them up when required, but they must be learned by heart, and in both directions; that is, an expression similar to the left side of a trigonometric formula must immediately recall the right side, and an expression similar to the right side must immediately recall the left side of the formula.

Trigonometric functions are defined on the circle, and on the right triangle.

Let in the circle, Fig. 28, the direction to the right and upward be considered as positive, to the left and downward as negative, and the angle α be counted from the positive horizontal $\overline{OA}$, counterclockwise as positive, clockwise as negative.

The *projector* s of the angle α , divided by the radius, is called $\sin \alpha$; the *projection* c of the angle α , divided by the radius, is called $\cos \alpha$.

The intercept t on the vertical tangent at the origin A , divided by the radius, is called $\tan \alpha$; the intercept ct on the horizontal tangent at B , or 90 deg., behind A , divided by the radius, is called $\cot \alpha$.

Thus, in Fig. 28,

$$\left. \begin{array}{ll} \sin \alpha = \frac{s}{r}; & \cos \alpha = \frac{c}{r}; \\ \tan \alpha = \frac{t}{r}; & \cot \alpha = \frac{ct}{r}. \end{array} \right\} \dots \dots \dots (1)$$

In the right triangle, Fig. 29, with the angles α and β , opposite respectively to the cathetes a and b , and with the hypotenuse c , the trigonometric functions are:

$$\left. \begin{aligned} \sin \alpha &= \cos \beta = \frac{a}{c}; & \cos \alpha &= \sin \beta = \frac{b}{c} \\ \tan \alpha &= \cot \beta = \frac{a}{b}; & \cot \alpha &= \tan \beta = \frac{b}{a} \end{aligned} \right\} \quad \dots \quad (2)$$

By the right triangle, only functions of angles up to 90 deg., or $\frac{\pi}{2}$, can be defined, while by the circle the trigonometric functions of any angle are given. Both representations thus must be so familiar to the engineer that he can see the trigo-

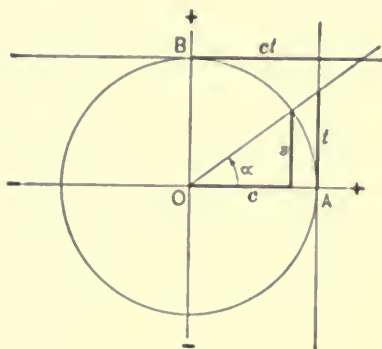


FIG. 28. Circular Trigonometric Functions.

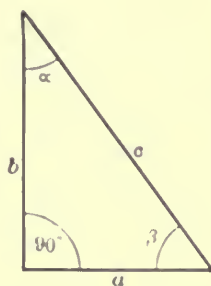


FIG. 29. Triangular Trigonometric Functions.

nometric functions and their variations with a change of the angle, and in most cases their numerical values, from the mental picture of the diagram.

67. Signs of Functions. In the first quadrant, Fig. 28, all trigonometric functions are positive.

In the second quadrant, Fig. 30, the $\sin \alpha$ is still positive, as s is in the upward direction, but $\cos \alpha$ is negative, since c is toward the left, and $\tan \alpha$ and $\cot \alpha$ also are negative, as t is downward, and ct toward the left.

In the third quadrant, Fig. 31, $\sin \alpha$ and $\cos \alpha$ are both

negative: s being downward, c toward the left; but $\tan \alpha$ and $\cot \alpha$ are again positive, as seen from t and ct in Fig. 31.

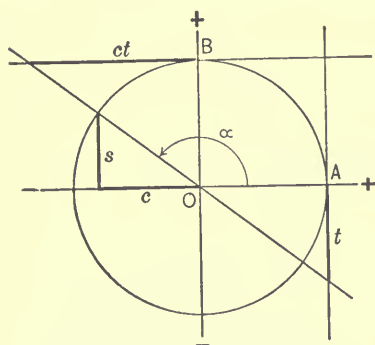


FIG. 30. Second Quadrant.

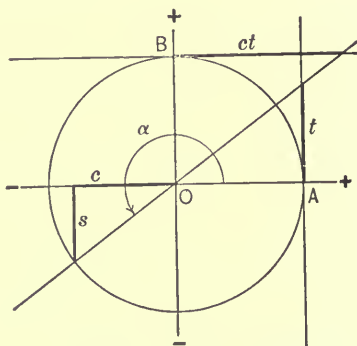


FIG. 31. Third Quadrant.

In the fourth quadrant, Fig. 32, $\sin \alpha$ is negative, as s is downward, but $\cos \alpha$ is again positive, as c is toward the right;

$\tan \alpha$ and $\cot \alpha$ are both negative, as seen from t and ct in Fig. 32.

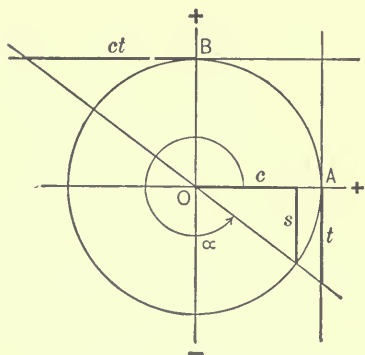


FIG. 32. Fourth Quadrant.

In the fifth quadrant all the trigonometric functions again have the same values as in the first quadrant, Fig. 28, that is, 360° , or 2π , or a multiple thereof, can be added to, or subtracted from the angle α , without changing the trigonometric functions, but these functions repeat after every 360° , or 2π ;

that is, have 2π or 360° as their period.

SIGNS OF FUNCTIONS

Function.	Positive.	Negative.
$\sin \alpha$	1st and 2d	3d and 4th quadrant
$\cos \alpha$	1st and 4th	2d and 3d "
$\tan \alpha$	1st and 3d	2d and 4th "
$\cot \alpha$	1st and 3d	2d and 4th "

(3)

68. Relations between $\sin \alpha$ and $\cos \alpha$. Between $\sin \alpha$ and $\cos \alpha$ the relation,

$$\sin^2 \alpha + \cos^2 \alpha = 1, \quad (4)$$

exists; hence,

$$\left. \begin{aligned} \sin \alpha &= \sqrt{1 - \cos^2 \alpha}; \\ \cos \alpha &= \sqrt{1 - \sin^2 \alpha}. \end{aligned} \right\} (4a)$$

Equation (4) is one of those which is frequently used in both directions. For instance, 1 may be substituted for the sum of the squares of $\sin \alpha$ and $\cos \alpha$, while in other cases $\sin^2 \alpha + \cos^2 \alpha$ may be substituted for 1. For instance,

$$\frac{1}{\cos^2 \alpha} = \frac{\sin^2 \alpha + \cos^2 \alpha}{\cos^2 \alpha} = \left(\frac{\sin \alpha}{\cos \alpha} \right)^2 + 1 = \tan^2 \alpha + 1.$$

Relations between Sines and Tangents.

$$\left. \begin{aligned} \tan \alpha &= \frac{\sin \alpha}{\cos \alpha}; \\ \cot \alpha &= \frac{\cos \alpha}{\sin \alpha}; \end{aligned} \right\} (5)$$

hence

$$\left. \begin{aligned} \cot \alpha &= \frac{1}{\tan \alpha}; \\ \tan \alpha &= \frac{1}{\cot \alpha}. \end{aligned} \right\} (5a)$$

As $\tan \alpha$ and $\cot \alpha$ are far less convenient for trigonometric calculations than $\sin \alpha$ and $\cos \alpha$, and therefore are less frequently applied in trigonometric calculations, it is not necessary to memorize the trigonometric formulas pertaining to $\tan \alpha$ and $\cot \alpha$, but where these functions occur, $\sin \alpha$ and $\cos \alpha$ are substituted for them by equations (5), and the calculations carried out with the latter functions, and $\tan \alpha$ or $\cot \alpha$ resubstituted in the final result, if the latter contains $\frac{\sin \alpha}{\cos \alpha}$, or its reciprocal.

In electrical engineering $\tan \alpha$ or $\cot \alpha$ frequently appears as the starting-point of calculation of the phase of alternating currents. For instance, if α is the phase angle of a vector

quantity, $\tan \alpha$ is given as the ratio of the vertical component over the horizontal component, or of the reactive component over the power component.

In this case, if

$$\tan \alpha = \frac{a}{b},$$

$$\sin \alpha = \frac{a}{\sqrt{a^2 + b^2}}, \quad \text{and} \quad \cos \alpha = \frac{b}{\sqrt{a^2 + b^2}}; \quad (5b)$$

or, if

$$\cot \alpha = \frac{c}{d},$$

$$\sin \alpha = \frac{d}{\sqrt{c^2 + d^2}}, \quad \text{and} \quad \cos \alpha = \frac{c}{\sqrt{c^2 + d^2}}. \quad (5c)$$

The secant functions, and versed sine functions are so little used in engineering, that they are of interest only as curiosities. They are defined by the following equations:

$$\sec \alpha = \frac{1}{\cos \alpha},$$

$$\operatorname{cosec} \alpha = \frac{1}{\sin \alpha},$$

$$\sin \operatorname{vers} \alpha = 1 - \sin \alpha,$$

$$\cos \operatorname{vers} \alpha = 1 - \cos \alpha.$$

69. Negative Angles. From the circle diagram of the trigonometric functions follows, as shown in Fig. 33, that when changing from a positive angle, that is, counterclockwise rotation, to a negative angle, that is, clockwise rotation, s , t , and ct reverse their direction, but c remains the same; that is,

$$\left. \begin{aligned} \sin (-\alpha) &= -\sin \alpha, \\ \cos (-\alpha) &= +\cos \alpha, \\ \tan (-\alpha) &= -\tan \alpha, \\ \cot (-\alpha) &= -\cot \alpha, \end{aligned} \right\} \quad (6)$$

$\cos \alpha$ thus is an "even function," while the three others are "odd functions."

Supplementary Angles. From the circle diagram of the trigonometric functions follows, as shown in Fig. 34, that by changing from an angle to its supplementary angle, s remains in the same direction, but c , t , and ct reverse their direction, and all four quantities retain the same numerical values, thus,

$$\left. \begin{aligned} \sin (\pi-\alpha) &= +\sin \alpha, \\ \cos (\pi-\alpha) &= -\cos \alpha, \\ \tan (\pi-\alpha) &= -\tan \alpha, \\ \cot (\pi-\alpha) &= -\cot \alpha. \end{aligned} \right\} \dots \dots \dots (7)$$

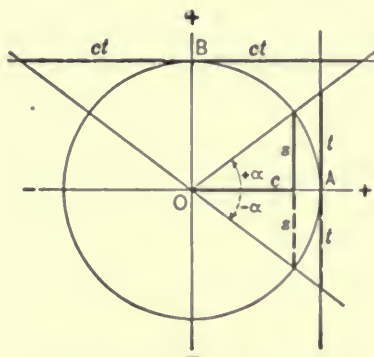


FIG. 33. Functions of Negative Angles.

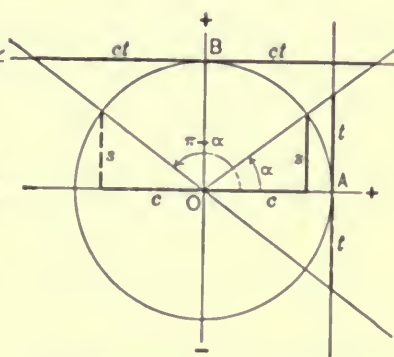


FIG. 34. Functions of Supplementary Angles.

Complementary Angles. Changing from an angle α to its complementary angle $90^\circ - \alpha$, or $\frac{\pi}{2} - \alpha$, as seen from Fig. 35, the signs remain the same, but s and c , and also t and ct exchange their numerical values, thus,

$$\left. \begin{aligned} \sin \left(\frac{\pi}{2} - \alpha \right) &= \cos \alpha, \\ \cos \left(\frac{\pi}{2} - \alpha \right) &= \sin \alpha, \\ \tan \left(\frac{\pi}{2} - \alpha \right) &= \cot \alpha, \\ \cot \left(\frac{\pi}{2} - \alpha \right) &= \tan \alpha. \end{aligned} \right\} \dots \dots \dots (8)$$

70. Angle $(\alpha \pm \pi)$. Adding, or subtracting π to an angle α , gives the same numerical values of the trigonometric functions

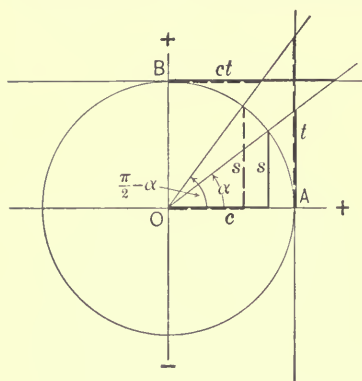


FIG. 35. Functions of Complementary Angles.

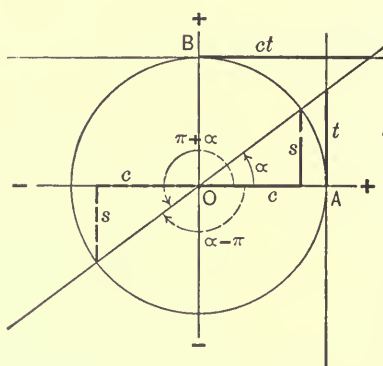


FIG. 36. Functions of Angles Plus or Minus π .

as α , as seen in Fig. 36, but the direction of s and c is reversed, while t and ct remain in the same direction, thus,

$$\left. \begin{aligned} \sin(\alpha \pm \pi) &= -\sin \alpha, \\ \cos(\alpha \pm \pi) &= -\cos \alpha, \\ \tan(\alpha \pm \pi) &= +\tan \alpha, \\ \cot(\alpha \pm \pi) &= +\cot \alpha. \end{aligned} \right\} \dots \dots \dots (9)$$

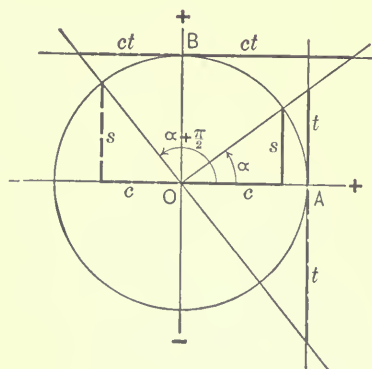


FIG. 37. Functions of Angles $+\frac{\pi}{2}$.

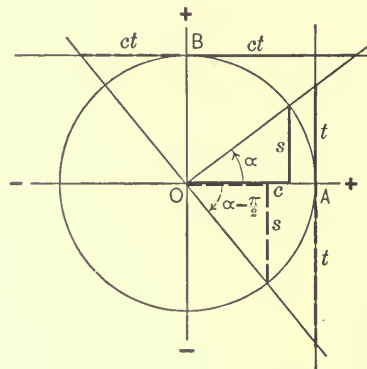


FIG. 38. Functions of Angles $-\frac{\pi}{2}$.

Angle $(\alpha \pm \frac{\pi}{2})$. Adding $\frac{\pi}{2}$, or 90 deg. to an angle α , interchanges the functions, s and c , and t and ct , and also reverses

the direction of the cosine, tangent, and cotangent, but leaves the sine in the same direction, since the sine is positive in the second quadrant, as seen in Fig. 37.

Subtracting $\frac{\pi}{2}$, or 90 deg. from angle α , interchanges the functions, s and c , and t and ct , and also reverses the direction, except that of the cosine, which remains in the same direction; that is, of the same sign, as the cosine is positive in the first and fourth quadrant, as seen in Fig. 38. Therefore,

$$\left. \begin{aligned} \sin \left(\alpha + \frac{\pi}{2} \right) &= +\cos \alpha, \\ \cos \left(\alpha + \frac{\pi}{2} \right) &= -\sin \alpha, \\ \tan \left(\alpha + \frac{\pi}{2} \right) &= -\cot \alpha, \\ \cot \left(\alpha + \frac{\pi}{2} \right) &= -\tan \alpha, \end{aligned} \right\} \dots \dots \dots (10)$$

$$\left. \begin{aligned} \sin \left(\alpha - \frac{\pi}{2} \right) &= -\cos \alpha, \\ \cos \left(\alpha - \frac{\pi}{2} \right) &= +\sin \alpha, \\ \tan \left(\alpha - \frac{\pi}{2} \right) &= -\cot \alpha, \\ \cot \left(\alpha - \frac{\pi}{2} \right) &= -\tan \alpha. \end{aligned} \right\} \dots \dots \dots (11)$$

Numerical Values. From the circle diagram, Fig. 28, etc., follows the numerical values:

$$\left. \begin{array}{cc|cc} \sin 0^\circ = 0 & \cos 0^\circ = 1 & \tan 0^\circ = 0 & \cot 0^\circ = \infty \\ \sin 30^\circ = \frac{1}{2} & \cos 30^\circ = \frac{1}{2} \sqrt{3} & \tan 45^\circ = 1 & \cot 45^\circ = 1 \\ \sin 45^\circ = \frac{1}{2} \sqrt{2} & \cos 45^\circ = \frac{1}{2} \sqrt{2} & \tan 90^\circ = \infty & \cot 90^\circ = 0 \\ \sin 60^\circ = \frac{1}{2} \sqrt{3} & \cos 60^\circ = \frac{1}{2} & \tan 135^\circ = -1 & \cot 135^\circ = -1 \\ \sin 90^\circ = 1 & \cos 90^\circ = 0 & \text{etc.} & \text{etc.} \\ \sin 120^\circ = \frac{1}{2} \sqrt{3} & \cos 120^\circ = -\frac{1}{2} & & \\ \text{etc.} & \text{etc.} & & \end{array} \right\} \dots (12)$$

71. Relations between Two Angles. The following relations are developed in text-books of trigonometry:

$$\left. \begin{aligned} \sin (\alpha+\beta) &= \sin \alpha \cos \beta + \cos \alpha \sin \beta, \\ \sin (\alpha-\beta) &= \sin \alpha \cos \beta - \cos \alpha \sin \beta, \\ \cos (\alpha+\beta) &= \cos \alpha \cos \beta - \sin \alpha \sin \beta, \\ \cos (\alpha-\beta) &= \cos \alpha \cos \beta + \sin \alpha \sin \beta, \end{aligned} \right\} \quad . \quad . \quad (13)$$

Herefrom follows, by combining these equations (13) in pairs:

$$\left. \begin{aligned} \cos \alpha \cos \beta &= \frac{1}{2} \{ \cos (\alpha+\beta) + \cos (\alpha-\beta) \}, \\ \sin \alpha \sin \beta &= \frac{1}{2} \{ \cos (\alpha-\beta) - \cos (\alpha+\beta) \}, \\ \sin \alpha \cos \beta &= \frac{1}{2} \{ \sin (\alpha+\beta) + \sin (\alpha-\beta) \}, \\ \cos \alpha \sin \beta &= \frac{1}{2} \{ \sin (\alpha+\beta) - \sin (\alpha-\beta) \}. \end{aligned} \right\} \quad . \quad . \quad (14)$$

By substituting α_1 for $(\alpha+\beta)$, and β_1 for $(\alpha-\beta)$ in these equations (14), gives the equations,

$$\left. \begin{aligned} \sin \alpha_1 + \sin \beta_1 &= 2 \sin \frac{\alpha_1 + \beta_1}{2} \cos \frac{\alpha_1 - \beta_1}{2}, \\ \sin \alpha_1 - \sin \beta_1 &= 2 \sin \frac{\alpha_1 - \beta_1}{2} \cos \frac{\alpha_1 + \beta_1}{2}, \\ \cos \alpha_1 + \cos \beta_1 &= 2 \cos \frac{\alpha_1 + \beta_1}{2} \cos \frac{\alpha_1 - \beta_1}{2}, \\ \cos \alpha_1 - \cos \beta_1 &= -2 \sin \frac{\alpha_1 + \beta_1}{2} \sin \frac{\alpha_1 - \beta_1}{2}. \end{aligned} \right\} \quad (15)$$

These three sets of equations are the most important trigonometric formulas. Their memorizing can be facilitated by noting that cosine functions lead to products of equal functions, sine functions to products of unequal functions, and inversely, products of equal functions resolve into cosine, products of unequal functions into sine functions. Also cosine functions show a reversal of the sign, thus: the cosine of a sum is given by a difference of products, the cosine of a difference by a sum, for the reason that with increasing angle the cosine function decreases, and the cosine of a sum of angles thus would be less than the cosine of the single angle.

Double Angles. From (13) follows, by substituting α for β :

$$\left. \begin{aligned} \sin 2\alpha &= 2 \sin \alpha \cos \alpha, \\ \cos 2\alpha &= \cos^2 \alpha - \sin^2 \alpha, \\ &= 2 \cos^2 \alpha - 1, \\ &= 1 - 2 \sin^2 \alpha. \end{aligned} \right\} \quad (16)$$

Herefrom follow

$$\sin^2 \alpha = \frac{1 - \cos 2\alpha}{2} \quad \text{and} \quad \cos^2 \alpha = \frac{1 + \cos 2\alpha}{2} \quad \left. \right\} . (16a)$$

72. Differentiation.

$$\left. \begin{aligned} \frac{d}{d\alpha} (\sin \alpha) &= +\cos \alpha, \\ \frac{d}{d\alpha} (\cos \alpha) &= -\sin \alpha. \end{aligned} \right\} \quad (17)$$

The sign of the latter differential is negative, as with an increase of angle α , the $\cos \alpha$ decreases.

Integration.

$$\left. \begin{aligned} \int \sin \alpha d\alpha &= -\cos \alpha, \\ \int \cos \alpha d\alpha &= +\sin \alpha. \end{aligned} \right\} \quad (18)$$

Herefrom follow the definite integrals:

$$\left. \begin{aligned} \int_c^{c+2\pi} \sin (\alpha + a) d\alpha &= 0; \\ \int_c^{c+2\pi} \cos (\alpha + a) d\alpha &= 0; \end{aligned} \right\} \quad (18a)$$

$$\left. \begin{aligned} \int_c^{c+\pi} \sin (\alpha + a) d\alpha &= 2 \cos (c + a); \\ \int_c^{c+\pi} \cos (\alpha + a) d\alpha &= -2 \sin (c + a); \end{aligned} \right\} \quad . . (18b)$$

$$\left. \begin{aligned} \int_{-\frac{\pi}{2}}^{+\frac{\pi}{2}} \sin \alpha d\alpha &= 0; \\ \int_0^{\pi} \cos \alpha d\alpha &= 0; \end{aligned} \right\} \dots \dots \dots (18c)$$

$$\left. \begin{aligned} \int_0^{\frac{\pi}{2}} \sin \alpha d\alpha &= +1; \\ \int_0^{\frac{\pi}{2}} \cos \alpha d\alpha &= +1. \end{aligned} \right\} \dots \dots \dots (18d)$$

73. Binomial. One of the most frequent trigonometric operations in electrical engineering is the transformation of the binomial, $a \cos \alpha + b \sin \alpha$, into a single trigonometric function, by the substitution, $a = c \cos p$ and $b = c \sin p$; hence,

$$a \cos \alpha + b \sin \alpha = c \cos (\alpha - p), \quad \dots \dots (19)$$

where

$$c = \sqrt{a^2 + b^2} \quad \text{and} \quad \tan p = \frac{b}{a}; \quad \dots \dots (20)$$

or, by the transformation, $a = c \sin q$ and $b = c \cos q$,

$$a \cos \alpha + b \sin \alpha = c \sin (\alpha + q), \quad \dots \dots (21)$$

where

$$c = \sqrt{a^2 + b^2} \quad \text{and} \quad \tan q = \frac{a}{b}. \quad \dots \dots (22)$$

74. Polyphase Relations.

$$\left. \begin{aligned} \sum_1^n i \cos \left(\alpha + a \pm \frac{2mi\pi}{n} \right) &= 0, \\ \sum_1^n i \sin \left(\alpha + a \pm \frac{2mi\pi}{n} \right) &= 0, \end{aligned} \right\} \dots \dots (23)$$

where m and n are integer numbers.

Proof. The points on the circle which defines the trigonometric function, by Fig. 28, of the angles $\left(\alpha + a \pm \frac{2mi\pi}{n} \right)$,

are corners of a regular polygon, inscribed in the circle and therefore having the center of the circle as center of gravity. For instance, for $n=7$, $m=2$, they are shown as $P_1, P_2, \dots, P_7$, in Fig. 39. The cosines of these angles are the projections on the vertical, the sines, the projections on the horizontal diameter, and as the sum of the projections of the corners of any polygon,

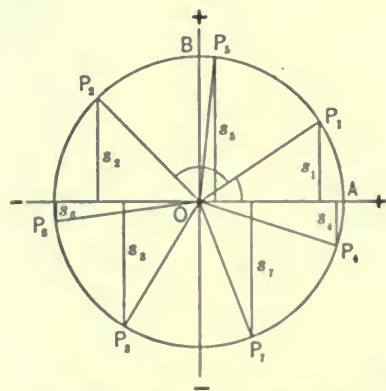


FIG. 39. Polyphase Relations.

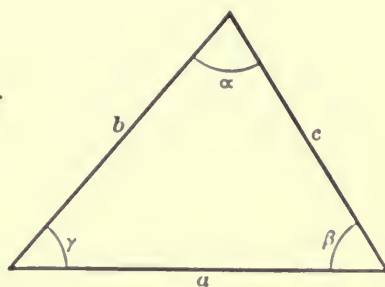


FIG. 40. Triangle.

on any line going through its center of gravity, is zero, both sums of equation (23) are zero.

$$\left. \begin{aligned} \sum_1^n i \cos \left(\alpha + a \pm \frac{2mi\pi}{n} \right) \cos \left(\alpha + b \pm \frac{2mi\pi}{n} \right) &= \frac{n}{2} \cos (a - b), \\ \sum_1^n i \sin \left(\alpha + a \pm \frac{2mi\pi}{n} \right) \sin \left(\alpha + b \pm \frac{2mi\pi}{n} \right) &= \frac{n}{2} \cos (a - b), \\ \sum_1^n i \sin \left(\alpha + a \pm \frac{2mi\pi}{n} \right) \cos \left(\alpha + b \pm \frac{2mi\pi}{n} \right) &= \frac{n}{2} \sin (a - b). \end{aligned} \right\} \quad (24)$$

These equations are proven by substituting for the products the single functions by equations (14), and substituting them in equations (23).

75. Triangle. If in a triangle α , β , and γ are the angles, opposite respectively to the sides a , b , c , Fig. 40, then,

$$\sin \alpha \div \sin \beta \div \sin \gamma = a \div b \div c, \quad . \quad . \quad . \quad (25)$$

$$\left. \begin{aligned} \cos \gamma &= \frac{a^2 + b^2 - c^2}{2ab}; \\ \text{or} \quad c^2 &= a^2 + b^2 - 2ab \cos \gamma. \end{aligned} \right\} \dots \dots \dots (26)$$

$$\left. \begin{aligned} \text{Area} &= \frac{ab \sin \gamma}{2} \\ &= \frac{c^2 \sin \alpha \sin \beta}{2 \sin \gamma}. \end{aligned} \right\} \dots \dots \dots (27)$$

B. TRIGONOMETRIC SERIES.

76. Engineering phenomena usually are either constant, transient, or periodic. Constant, for instance, is the terminal voltage of a storage-battery and the current taken from it through a constant resistance. Transient phenomena occur during a change in the condition of an electric circuit, as a change of load; or, disturbances entering the circuit from the outside or originating in it, etc. Periodic phenomena are the alternating currents and voltages, pulsating currents as those produced by rectifiers, the distribution of the magnetic flux in the air-gap of a machine, or the distribution of voltage around the commutator of the direct-current machine, the motion of the piston in the steam-engine cylinder, the variation of the mean daily temperature with the seasons of the year, etc.

The characteristic of a periodic function, $y=f(x)$, is, that at constant intervals of the independent variable x , called cycles or periods, the same values of the dependent variable y occur.

Most periodic functions of engineering are functions of time or of space, and as such have the characteristic of univalence; that is, to any value of the independent variable x can correspond only one value of the dependent variable y . In other words, at any given time and given point of space, any physical phenomenon can have one numerical value only, and therefore must be represented by a univalent function of time and space.

Any univalent periodic function,

$$y=f(x), \quad \dots \dots \dots (1)$$

can be expressed by an infinite trigonometric series, or Fourier series, of the form,

$$y = a_0 + a_1 \cos cx + a_2 \cos 2cx + a_3 \cos 3cx + \dots \\ + b_1 \sin cx + b_2 \sin 2cx + b_3 \sin 3cx + \dots \quad (2)$$

or, substituting for convenience, $cx = \theta$, this gives

$$y = a_0 + a_1 \cos \theta + a_2 \cos 2\theta + a_3 \cos 3\theta + \dots \\ + b_1 \sin \theta + b_2 \sin 2\theta + b_3 \sin 3\theta + \dots \quad (3)$$

or, combining the sine and cosine functions by the binomial (par. 73),

$$y = a_0 + c_1 \cos (\theta - \beta_1) + c_2 \cos (2\theta - \beta_2) + c_3 \cos (3\theta - \beta_3) + \dots \\ = a_0 + c_1 \sin (\theta + \gamma_1) + c_2 \sin (2\theta + \gamma_2) + c_3 \sin (3\theta + \gamma_3) + \dots \quad (4)$$

where

$$\left. \begin{aligned} c_n &= \sqrt{a_n^2 + b_n^2}; \\ \tan \beta_n &= \frac{b_n}{a_n}; \\ \tan \gamma_n &= \frac{a_n}{b_n}. \end{aligned} \right\} \quad (5)$$

or

The proof hereof is given by showing that the coefficients a_n and b_n of the series (3) can be determined from the numerical values of the periodic function (1), thus,

$$y = f(x) = f_0(\theta). \quad (6)$$

Since, however, the trigonometric function, and therefore also the series of trigonometric functions (3) is univalent, it follows that the periodic function (6), $y = f_0(\theta)$, must be univalent, to be represented by a trigonometric series.

77. The most important periodic functions in electrical engineering are the alternating currents and e.m.fs. Usually they are, in first approximation, represented by a single trigonometric function, as:

$$i = i_0 \cos (\theta - \omega);$$

or,

$$e = e_0 \sin (\theta - \delta);$$

that is, they are assumed as sine waves.

Theoretically, obviously this condition can never be perfectly attained, and frequently the deviation from sine shape is sufficient to require practical consideration, especially in those cases, where the electric circuit contains electrostatic capacity, as is for instance, the case with long-distance transmission lines, underground cable systems, high potential transformers, etc.

However, no matter how much the alternating or other periodic wave differs from simple sine shape—that is, however much the wave is “distorted,” it can always be represented by the trigonometric series (3).

As illustration the following applications of the trigonometric series to engineering problems may be considered:

(A) The determination of the equation of the periodic function; that is, the evolution of the constants a_n and b_n of the trigonometric series, if the numerical values of the periodic function are given. Thus, for instance, the wave of an alternator may be taken by oscillograph or wave-meter, and by measuring from the oscillograph, the numerical values of the periodic function are derived for every 10 degrees, or every 5 degrees, or every degree, depending on the accuracy required. The problem then is, from the numerical values of the wave, to determine its equation. While the oscillograph shows the shape of the wave, it obviously is not possible therefrom to calculate other quantities, as from the voltage the current under given circuit conditions, if the wave shape is not first represented by a mathematical expression. It therefore is of importance in engineering to translate the picture or the table of numerical values of a periodic function into a mathematical expression thereof.

(B) If one of the engineering quantities, as the e.m.f. of an alternator or the magnetic flux in the air-gap of an electric machine, is given as a general periodic function in the form of a trigonometric series, to determine therefrom other engineering quantities, as the current, the generated e.m.f., etc.

A. Evaluation of the Constants of the Trigonometric Series from the Instantaneous Values of the Periodic Function.

78. Assuming that the numerical values of a univalent periodic function $y=f_0(\theta)$ are given; that is, for every value of θ , the corresponding value of y is known, either by graphical representation, Fig. 41; or, in tabulated form, Table I, but

the equation of the periodic function is not known. It can be represented in the form,

$$y = a_0 + a_1 \cos \theta + a_2 \cos 2\theta + a_3 \cos 3\theta + \dots + a_n \cos n\theta + \dots \\ + b_1 \sin \theta + b_2 \sin 2\theta + b_3 \sin 3\theta + \dots + b_n \sin n\theta + \dots, \quad (7)$$

and the problem now is, to determine the coefficients $a_0, a_1, a_2 \dots b_1, b_2 \dots$.

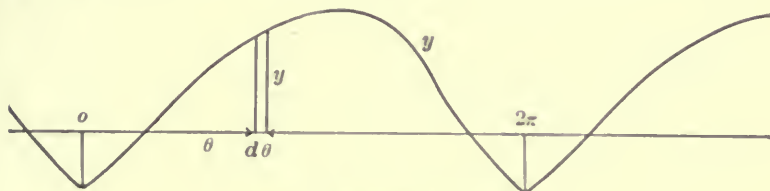


FIG. 41. Periodic Functions.

TABLE I.

θ	y	θ	y	θ	y	θ	y
0	-60	90	+ 50	180	+ 122	270	+ 85
10	-49	100	+ 61	190	+ 124	280	+ 65
20	-38	110	+ 71	200	+ 126	290	+ 35
30	-26	120	+ 81	210	+ 125	300	+ 17
40	-12	130	+ 90	220	+ 123	310	0
50	0	140	+ 99	230	+ 120	320	- 13
60	+ 11	150	+ 107	240	+ 116	330	- 26
70	+ 27	160	+ 114	250	+ 110	340	- 38
80	+ 39	170	+ 119	260	+ 100	350	- 49
90	+ 50	180	+ 122	270	+ 85	360	- 60

Integrate the equation (7) between the limits 0 and 2π :

$$\int_0^{2\pi} y d\theta = a_0 \int_0^{2\pi} d\theta + a_1 \int_0^{2\pi} \cos \theta d\theta + a_2 \int_0^{2\pi} \cos 2\theta d\theta + \dots \\ + a_n \int_0^{2\pi} \cos n\theta d\theta + \dots + b_1 \int_0^{2\pi} \sin \theta d\theta + \\ + b_2 \int_0^{2\pi} \sin 2\theta d\theta + \dots + b_n \int_0^{2\pi} \sin n\theta d\theta + \dots \\ = a_0 \left/ \theta \right/_0^{2\pi} + a_1 \left/ \sin \theta \right/_0^{2\pi} + a_2 \left/ \frac{\sin 2\theta}{2} \right/_0^{2\pi} + \dots \\ + a_n \left/ \frac{\sin n\theta}{n} \right/_0^{2\pi} + \dots - b_1 \left/ \cos \theta \right/_0^{2\pi} \\ - b_2 \left/ \frac{\cos 2\theta}{2} \right/_0^{2\pi} - \dots - b_n \left/ \frac{\cos n\theta}{n} \right/_0^{2\pi} + \dots$$

All the integrals containing trigonometric functions vanish, as the trigonometric function has the same value at the upper limit 2π as at the lower limit 0, that is,

$$\left/ \frac{\cos n\theta}{n} \right/_0^{2\pi} = \frac{1}{n}(\cos 2n\pi - \cos 0) = 0;$$

$$\left/ \frac{\sin n\theta}{n} \right/_0^{2\pi} = \frac{1}{n}(\sin 2n\pi - \sin 0) = 0,$$

and the result is

$$\int_0^{2\pi} y d\theta = a_0 \left/ \theta \right/_0^{2\pi} = 2\pi a_0;$$

hence

$$a_0 = \frac{1}{2\pi} \int_0^{2\pi} y d\theta. \quad . \quad . \quad . \quad . \quad . \quad (8)$$

$y d\theta$ is an element of the area of the curve y , Fig. 41, and $\int_0^{2\pi} y d\theta$ thus is the area of the periodic function y , for one period; that is,

$$a_0 = \frac{1}{2\pi} A, \quad . \quad . \quad . \quad . \quad . \quad (9)$$

where A = area of the periodic function $y = f_0(\theta)$, for one period; that is, from $\theta = 0$ to $\theta = 2\pi$.

2π is the horizontal width of this area A , and $\frac{A}{2\pi}$ thus is the area divided by the width of it; that is, it is the average height of the area A of the periodic function y ; or, in other words, it is the average value of y . Therefore,

$$a_0 = \text{avg. } (y)_0^{2\pi}. \quad . \quad . \quad . \quad . \quad . \quad (10)$$

The first coefficient, a_0 , thus, is the average value of the instantaneous values of the periodic function y , between $\theta = 0$ and $\theta = 2\pi$.

Therefore, averaging the values of y in Table I, gives the first constant a_0 .

79. To determine the coefficient a_n , multiply equation (7) by $\cos n\theta$, and then integrate from 0 to 2π , for the purpose of making the trigonometric functions vanish. This gives

$$\begin{aligned}
\int_0^{2\pi} y \cos n\theta d\theta &= a_0 \int_0^{2\pi} \cos n\theta d\theta + a_1 \int_0^{2\pi} \cos n\theta \cos \theta d\theta + \\
&+ a_2 \int_0^{2\pi} \cos n\theta \cos 2\theta d\theta + \dots + a_n \int_0^{2\pi} \cos^2 n\theta d\theta + \dots \\
&+ b_1 \int_0^{2\pi} \cos n\theta \sin \theta d\theta + b_2 \int_0^{2\pi} \cos n\theta \sin 2\theta d\theta + \dots \\
&+ b_n \int_0^{2\pi} \cos n\theta \sin n\theta d\theta + \dots
\end{aligned}$$

Hence, by the trigonometric equations of the preceding section:

$$\begin{aligned}
\int_0^{2\pi} y \cos n\theta d\theta &= a_0 \int_0^{2\pi} \cos n\theta d\theta + a_1 \int_0^{2\pi} \frac{1}{2} [\cos (n+1)\theta + \cos (n-1)\theta] d\theta \\
&+ a_2 \int_0^{2\pi} \frac{1}{2} [\cos (n+2)\theta + \cos (n-2)\theta] d\theta + \dots \\
&+ a_n \int_0^{2\pi} \frac{1}{2} (1 + \cos 2n\theta) d\theta + \dots \\
&+ b_1 \int_0^{2\pi} \frac{1}{2} [\sin (n+1)\theta - \sin (n-1)\theta] d\theta \\
&+ b_2 \int_0^{2\pi} \frac{1}{2} [\sin (n+2)\theta - \sin (n-2)\theta] d\theta + \dots \\
&+ b_n \int_0^{2\pi} \frac{1}{2} \sin 2n\theta d\theta + \dots
\end{aligned}$$

All these integrals of trigonometric functions give trigonometric functions, and therefore vanish between the limits 0 and 2π , and there only remains the first term of the integral multiplied with a_n , which does not contain a trigonometric function, and thus remains finite:

$$a_n \int_0^{2\pi} \frac{1}{2} d\theta = a_n \left(\frac{\theta}{2} \right)_0^{2\pi} = a_n \pi,$$

and therefore,

$$\int_0^{2\pi} y \cos n\theta d\theta = a_n \pi;$$

hence

$$a_n = \frac{1}{\pi} \int_0^{2\pi} y \cos n\theta d\theta. \quad \dots \dots \dots (11)$$

If the instantaneous values of y are multiplied with $\cos n\theta$, and the product $y_n = y \cos n\theta$ plotted as a curve, $y \cos n\theta d\theta$ is an element of the area of this curve, shown for $n=3$ in Fig. 42, and thus $\int_0^{2\pi} y \cos n\theta d\theta$ is the area of this curve; that is,

$$a_n = \frac{1}{\pi} A_n, \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (12)$$

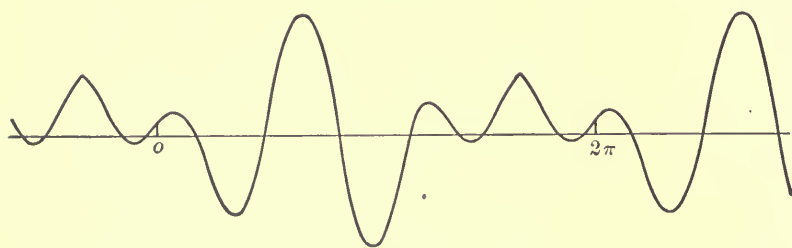


FIG. 42. Curve of $y \cos 3\theta$.

where A_n is the area of the curve $y \cos n\theta$, between $\theta=0$ and $\theta=2\pi$.

As 2π is the width of this area A_n , $\frac{A_n}{2\pi}$ is the average height of this area; that is, is the average value of $y \cos n\theta$, and $\frac{1}{\pi} A_n$ thus is twice the average value of $y \cos n\theta$; that is,

$$a_n = 2 \text{ avg. } (y \cos n\theta)_0^{2\pi}. \quad . \quad . \quad . \quad . \quad . \quad . \quad (13)$$

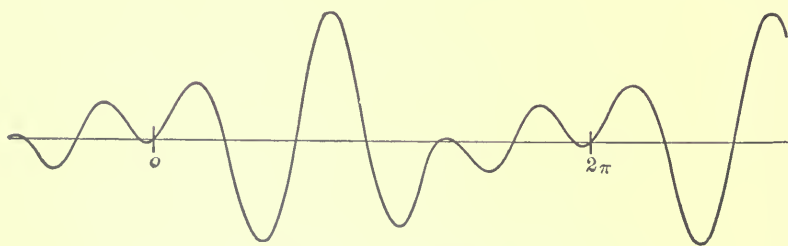


FIG. 43. Curve of $y \sin 3\theta$.

The coefficient a_n of $\cos n\theta$ is derived by multiplying all the instantaneous values of y by $\cos n\theta$, and taking twice the average of the instantaneous values of this product $y \cos n\theta$.

80. b_n is determined in the analogous manner by multiplying y by $\sin n\theta$ and integrating from 0 to 2π ; by the area of the curve $y \sin n\theta$, shown in Fig. 43, for $n=3$,

$$\begin{aligned}
 \int_0^{2\pi} y \sin n\theta d\theta &= a_0 \int_0^{2\pi} \sin n\theta d\theta + a_1 \int_0^{2\pi} \sin n\theta \cos \theta d\theta \\
 &+ a_2 \int_0^{2\pi} \sin n\theta \cos 2\theta d\theta + \dots + a_n \int_0^{2\pi} \sin n\theta \cos n\theta d\theta + \dots \\
 &+ b_1 \int_0^{2\pi} \sin n\theta \sin \theta d\theta + b_2 \int_0^{2\pi} \sin n\theta \sin 2\theta d\theta + \dots \\
 &+ b_n \int_0^{2\pi} \sin^2 n\theta d\theta + \dots \\
 &= a_0 \int_0^{2\pi} \sin n\theta d\theta + a_1 \int_0^{2\pi} \frac{1}{2} [\sin (n+1)\theta + \sin (n-1)\theta] d\theta \\
 &+ a_2 \int_0^{2\pi} \frac{1}{2} [\sin (n+2)\theta + \sin (n-2)\theta] d\theta + \dots \\
 &+ a_n \int_0^{2\pi} \frac{1}{2} \sin 2n\theta d\theta + \dots \\
 &+ b_1 \int_0^{2\pi} \frac{1}{2} [\cos (n-1)\theta - \cos (n+1)\theta] d\theta \\
 &+ b_2 \int_0^{2\pi} \frac{1}{2} [\cos (n-2)\theta - \cos (n+2)\theta] d\theta + \dots \\
 &+ b_n \int_0^{2\pi} \frac{1}{2} [1 - \cos 2n\theta] d\theta + \dots \\
 &= b_n \int_0^{2\pi} \frac{1}{2} d\theta = b_n \pi;
 \end{aligned}$$

hence,

$$b_n = \frac{1}{\pi} \int_0^{2\pi} y \sin n\theta d\theta \quad \dots \quad (14)$$

$$= \frac{1}{\pi} A_n', \quad \dots \quad (15)$$

where A_n' is the area of the curve $y_n' = y \sin n\theta$. Hence,

$$b_n = 2 \text{ avg. } (y \sin n\theta)_0^{2\pi}, \quad \dots \quad (16)$$

and the coefficient of $\sin n\theta$ thus is derived by multiplying the instantaneous values of y with $\sin n\theta$, and then averaging, as twice the average of $y \sin n\theta$.

81. Any univalent periodic function, of which the numerical values y are known, can thus be expressed numerically by the equation,

$$y = a_0 + a_1 \cos \theta + a_2 \cos 2\theta + \dots + a_n \cos n\theta + \dots$$

$$+ b_1 \sin \theta + b_2 \sin 2\theta + \dots + b_n \sin n\theta + \dots, \quad (17)$$

where the coefficients $a_0, a_1, a_2, \dots, b_1, b_2, \dots$, are calculated as the averages:

$$\left. \begin{aligned} a_0 &= \text{avg. } (y)_0^{2\pi}; \\ a_1 &= 2 \text{ avg. } (y \cos \theta)_0^{2\pi}; & b_1 &= 2 \text{ avg. } (y \sin \theta)_0^{2\pi}; \\ a_2 &= 2 \text{ avg. } (y \cos 2\theta)_0^{2\pi}; & b_2 &= 2 \text{ avg. } (y \sin 2\theta)_0^{2\pi}; \\ a_n &= 2 \text{ avg. } (y \cos n\theta)_0^{2\pi}; & b_n &= 2 \text{ avg. } (y \sin n\theta)_0^{2\pi}; \end{aligned} \right\} \quad (18)$$

Hereby any individual harmonic can be calculated, without calculating the preceding harmonics.

For instance, let the generator e.m.f. wave, Fig. 44, Table II, column 2, be impressed upon an underground cable system

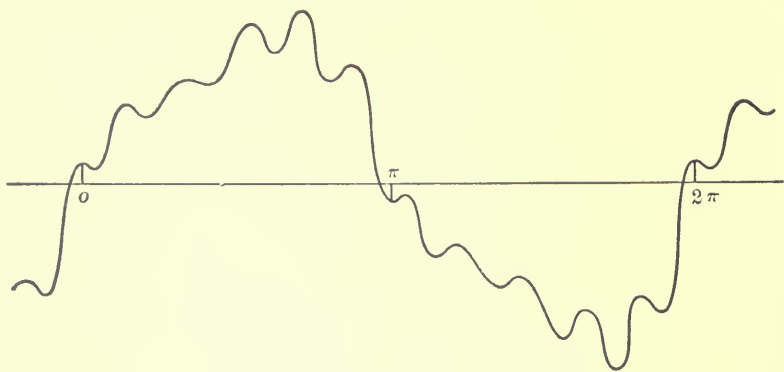


FIG. 44. Generator e.m.f. wave.

of such constants (capacity and inductance), that the natural frequency of the system is 670 cycles per second, while the generator frequency is 60 cycles. The natural frequency of the

circuit is then close to that of the 11th harmonic of the generator wave, 660 cycles, and if the generator voltage contains an appreciable 11th harmonic, trouble may result from a resonance rise of voltage of this frequency; therefore, the 11th harmonic of the generator wave is to be determined, that is, a_{11} and b_{11} calculated, but the other harmonics are of less importance.

TABLE II

θ	y	$\cos 11\theta$	$\sin 11\theta$	$y \cos 11\theta$	$y \sin 11\theta$
0	5	+1.000	0	+5.0	0
10	4	-0.342	+0.940	-1.4	+3.8
20	20	-0.766	-0.643	-15.3	-12.9
30	22	+0.866	-0.500	+19.1	-11.0
40	19	+0.174	+0.985	+3.3	+18.7
50	25	-0.985	-0.174	-24.6	-4.3
60	29	+0.500	-0.866	+14.5	-25.1
70	29	+0.643	+0.766	+18.6	+22.2
80	30	-0.940	+0.342	-28.2	+10.3
90	38	0	-1.000	0	-38.0
100	46	+0.940	+0.342	+43.3	+15.7
110	38	-0.643	+0.766	-24.4	+29.2
120	41	-0.500	-0.866	-20.5	-35.5
130	50	+0.985	-0.174	+49.2	-8.7
140	32	-0.174	+0.985	-5.6	+31.5
150	30	-0.866	-0.500	-26.0	-15.0
160	33	+0.766	-0.643	+25.3	-21.3
170	7	+0.342	+0.940	+2.2	-21.3
180	-5				
Total				+34.5	-29.8
Divided by 9				+3.83 = a_{11}	-3.31 = b_{11}

In the third column of Table II thus are given the values of $\cos 11\theta$, in the fourth column $\sin 11\theta$, in the fifth column $y \cos 11\theta$, and in the sixth column $y \sin 11\theta$. The former gives as average +1.915, hence $a_{11} = +3.83$, and the latter gives as average -1.655, hence $b_{11} = -3.31$, and the 11th harmonic of the generator wave is

$$\begin{aligned}
 a_{11} \cos 11\theta + b_{11} \sin 11\theta &= 3.83 \cos 11\theta - 3.31 \sin 11\theta \\
 &= 5.07 \cos (11\theta + 41^\circ),
 \end{aligned}$$

hence, its effective value is

$$\frac{5.07}{\sqrt{2}} = 3.58,$$

while the effective value of the total generator wave, that is, the square root of the mean squares of the instantaneous values y , is

$$e = 30.5,$$

thus the 11th harmonic is 11.8 per cent of the total voltage, and whether such a harmonic is safe or not, can now be determined from the circuit constants, more particularly its resistance.

82. In general, the successive harmonics decrease; that is, with increasing n , the values of a_n and b_n become smaller, and when calculating a_n and b_n by equation (18), for higher values of n they are derived as the small averages of a number of large quantities, and the calculation then becomes inconvenient and less correct.

Where the entire series of coefficients a_n and b_n is to be calculated, it thus is preferable not to use the complete periodic function y , but only the residual left after subtracting the harmonics which have already been calculated; that is, after a_0 has been calculated, it is subtracted from y , and the difference, $y_1 = y - a_0$, is used for the calculation of a_1 and b_1 .

Then $a_1 \cos \theta + b_1 \sin \theta$ is subtracted from y_1 , and the difference,

$$\begin{aligned} y_2 &= y_1 - (a_1 \cos \theta + b_1 \sin \theta) \\ &= y - (a_0 + a_1 \cos \theta + b_1 \sin \theta), \end{aligned}$$

is used for the calculation of a_2 and b_2 .

Then $a_2 \cos 2\theta + b_2 \sin 2\theta$ is subtracted from y_2 , and the rest, y_3 , used for the calculation of a_3 and b_3 , etc.

In this manner a higher accuracy is derived, and the calculation simplified by having the instantaneous values of the function of the same magnitude as the coefficients a_n and b_n .

As illustration, is given in Table III the calculation of the first three harmonics of the pulsating current, Fig. 41, Table I:

83. In electrical engineering, the most important periodic functions are the alternating currents and voltages. Due to the constructive features of alternating-current generators, alternating voltages and currents are almost always symmetrical waves; that is, the periodic function consists of alternate half-waves, which are the same in shape, but opposite in direction, or in other words, the instantaneous values from 180 deg. to 360 deg. are the same numerically, but opposite in sign, from the instantaneous values between 0 to 180 deg., and each cycle or period thus consists of two equal but opposite half cycles, as shown in Fig. 44. In the earlier days of electrical engineering, the frequency has for this reason frequently been expressed by the number of half-waves or alternations.

In a symmetrical wave, those harmonics which produce a difference in the shape of the positive and the negative half-wave, cannot exist; that is, their coefficients a and b must be zero. Only those harmonics can exist in which an increase of the angle θ by 180 deg., or π , reverses the sign of the function. This is the case with $\cos n\theta$ and $\sin n\theta$, if n is an odd number. If, however, n is an even number, an increase of θ by π increases the angle $n\theta$ by 2π or a multiple thereof, thus leaves $\cos n\theta$ and $\sin n\theta$ with the same sign. The same applies to a_0 . Therefore, symmetrical alternating waves comprise only the odd harmonics, but do not contain even harmonics or a constant term, and thus are represented by

$$y = a_1 \cos \theta + a_3 \cos 3\theta + a_5 \cos 5\theta + \dots \\ + b_1 \sin \theta + b_3 \sin 3\theta + b_5 \sin 5\theta + \dots \quad (19)$$

When calculating the coefficients a_n and b_n of a symmetrical wave by the expression (18), it is sufficient to average from 0 to π ; that is, over one half-wave only. In the second half-wave, $\cos n\theta$ and $\sin n\theta$ have the opposite sign as in the first half-wave, if n is an odd number, and since y also has the opposite sign in the second half-wave, $y \cos n\theta$ and $y \sin n\theta$ in the second half-wave traverses again the same values, with the same sign, as in the first half-wave, and their average thus is given by averaging over one half-wave only.

Therefore, a symmetrical univalent periodic function, as an

TABLE

θ	ν	$\nu_1 = \nu - a_0$	$\nu_1 \cos \theta$	$\nu_1 \sin \theta$	$c_1 = a_1 \cos \theta + b_1 \sin \theta$	$\nu_2 = \nu_0 - c_1$
0	-60	-111	-111	0	-84	-27
10	-49	-100	-98	-17	-85	-15
20	-38	-89	-84	-30	-83	-6
30	-26	-77	-67	-38	-79	+2
40	-12	-63	-48	-40	-72	9
50	0	-51	-33	-39	-63	12
60	+11	-40	-20	-35	-52	12
70	27	-24	-8	-23	-40	16
80	39	-12	-2	-12	-26	14
90	50	-1	0	-1	-11	10
100	61	+10	-2	+10	+4	6
110	71	20	-7	+19	18	+2
120	81	30	-15	+26	32	-2
130	90	39	-25	+30	45	-6
140	99	48	-37	+31	58	-10
150	107	56	-49	+28	67	-11
160	114	63	-59	+22	75	-12
170	119	68	-67	+12	81	-13
180	122	71	-71	0	84	-13
190	124	73	-72	-13	85	-12
200	126	75	-71	-26	83	-8
210	125	74	-64	-37	79	-5
220	123	72	-55	-47	72	0
230	120	69	-44	-53	63	+6
240	116	65	-32	-28	52	13
250	110	59	-20	-56	40	19
260	100	49	-9	-48	26	23
270	85	34	0	-34	11	23
280	65	+14	+2	-14	-4	18
290	35	-16	-5	+15	-18	+2
300	+17	-34	-17	+30	-32	-2
310	0	-51	-33	+39	-45	-6
320	-13	-64	-49	+41	-58	-6
330	-26	-75	-65	+37	-67	-8
340	-38	-89	-84	+30	-75	-14
350	-49	-100	-99	+17	-81	-19
Total . . +1826		Total -1520		-204		Total
Divided		Divided by				Divided by 18 . . .
by 36 ... + 50.7 = a_0		18 -84.4 = a_1		-11.3 = b_1		

III.

$\nu_2 \cos 2\theta$	$\nu_2 \sin 2\theta$	$c_2 = a_2 \cos 2\theta + b_2 \sin 2\theta$	$\nu_2 = \nu_1 - c_2$	$\nu_2 \cos 3\theta$	$\nu_2 \sin 3\theta$	θ
-27	0	-15	-12	-12	0	0
-14	-5	-12	-3	-3	-1	10
-5	-4	-7	+1	0	+1	20
+1	+2	-1	+3	0	+3	30
+2	+9	+4	+5	-2	+4	40
-2	+12	11	+1	-1	0	50
-6	+10	13	-1	+1	0	60
-12	+10	15	+1	-1	0	70
-13	+5	16	-2	+1	+2	80
-10	0	15	-5	0	+5	90
-6	-2	12	-6	-3	+5	100
-2	-1	7	-5	-4	+2	110
+1	+2	+1	-3	-3	0	120
+1	+6	-4	-2	-2	-1	130
-2	+10	-11	+1	0	+1	140
-5	+10	-13	+2	0	+2	150
-9	+8	-15	+3	-1	+3	160
-12	-4	-16	+3	-3	+1	170
-13	0	-15	+2	-2	0	180
-11	-4	-12	0	0	0	190
-6	-6	-7	-1	0	-1	200
-2	-4	-1	-4	0	-4	210
0	0	+4	-4	-2	-4	220
-1	+6	11	-5	-4	-2	230
-6	+11	13	0	0	0	240
-15	+12	15	+4	+4	+2	250
-22	+8	16	+7	+3	+6	260
-23	0	15	+8	0	+8	270
-17	-6	12	+6	-3	+5	280
-2	-1	7	-5	+4	-2	290
+1	+2	+1	-3	+3	0	300
+1	+6	-4	-2	+2	+1	310
-1	+6	-11	+5	2	4	320
-4	+7	-13	+5	0	5	330
-11	+9	-15	+1	0	-1	340
-18	+6	-16	3	-3	+1	350
-270	+120	Total	.	-33	+27	
-15.0 = a_2	+6.7 = b_2	Divided by 18	.	-1.8 = a_2	+1.5 = b_2	

alternating voltage and current usually is, can be represented by the expression,

$$y = a_1 \cos \theta + a_3 \cos 3 \theta + a_5 \cos 5 \theta + a_7 \cos 7 \theta + \dots \\ + b_1 \sin \theta + b_3 \sin 3 \theta + b_5 \sin 5 \theta + b_7 \sin 7 \theta + \dots; \quad (20)$$

where,

$$\left. \begin{aligned} a_1 &= 2 \text{ avg. } (y \cos \theta)_0^\pi; & b_1 &= 2 \text{ avg. } (y \sin \theta)_0^\pi; \\ a_3 &= 2 \text{ avg. } (y \cos 3\theta)_0^\pi; & b_3 &= 2 \text{ avg. } (y \sin 3\theta)_0^\pi; \\ a_5 &= 2 \text{ avg. } (y \cos 5\theta)_0^\pi; & b_5 &= 2 \text{ avg. } (y \sin 5\theta)_0^\pi; \\ a_7 &= 2 \text{ avg. } (y \cos 7\theta)_0^\pi; & b_7 &= 2 \text{ avg. } (y \sin 7\theta)_0^\pi. \end{aligned} \right\} \quad (21)$$

84. From 180 deg. to 360 deg., the even harmonics have the same, but the odd harmonics the opposite sign as from 0 to 180 deg. Therefore adding the numerical values in the range from 180 deg. to 360 deg. to those in the range from 0 to 180 deg., the odd harmonics cancel, and only the even harmonics remain. Inversely, by subtracting, the even harmonics cancel, and the odd ones remain.

Hereby the odd and the even harmonics can be separated. If $y = y(\theta)$ are the numerical values of a periodic function from 0 to 180 deg., and $y' = y(\theta + \pi)$ the numerical values of the same function from 180 deg. to 360 deg.,

$$y_2(\theta) = \frac{1}{2} \{ y(\theta) + y(\theta + \pi) \}, \quad \dots \quad (22)$$

is a periodic function containing only the even harmonics, and

$$y_1(\theta) = \frac{1}{2} \{ y(\theta) - y(\theta + \pi) \} \quad \dots \quad (23)$$

is a periodic function containing only the odd harmonics; that is:

$$y_1(\theta) = a_1 \cos \theta + a_3 \cos 3\theta + a_5 \cos 5\theta + \dots \\ + b_1 \sin \theta + b_3 \sin 3\theta + b_5 \sin 5\theta + \dots; \quad \dots \quad (24)$$

$$y_2(\theta) = a_2 \cos 2\theta + a_4 \cos 4\theta + \dots \\ + b_2 \sin 2\theta + b_4 \sin 4\theta + \dots \quad \dots \quad (25)$$

and the complete function is

$$y(\theta) = y_1(\theta) + y_2(\theta). \quad \dots \quad (26)$$

By this method it is convenient to determine whether even harmonics are present, and if they are present, to separate them from the odd harmonics.

Before separating the even harmonics and the odd harmonics, it is usually convenient to separate the constant term a_0 from the periodic function y , by averaging the instantaneous values of y from 0 to 360 deg. The average then gives a_0 , and subtracted from the instantaneous values of y , gives

$$y_0(\theta) = y(\theta) - a_0 \quad . \quad . \quad . \quad . \quad . \quad (27)$$

as the instantaneous values of the alternating component of the periodic function; that is, the component y_0 contains only the trigonometric functions, but not the constant term. y_0 is then resolved into the odd series y_1 , and the even series y_2 .

85. The alternating wave y_0 consists of the cosine components:

$$u(\theta) = a_1 \cos \theta + a_2 \cos 2\theta + a_3 \cos 3\theta + a_4 \cos 4\theta + \dots, \quad (28)$$

and the sine components:

$$v(\theta) = b_1 \sin \theta + b_2 \sin 2\theta + b_3 \sin 3\theta + b_4 \sin 4\theta + \dots; \quad (29)$$

that is,

$$y_0(\theta) = u(\theta) + v(\theta). \quad . \quad . \quad . \quad . \quad . \quad (30)$$

The cosine functions retain the same sign for negative angles $(-\theta)$, as for positive angles $(+\theta)$, while the sine functions reverse their sign; that is,

$$u(-\theta) = +u(\theta) \quad \text{and} \quad v(-\theta) = -v(\theta). \quad . \quad . \quad . \quad . \quad (31)$$

Therefore, if the values of y_0 for positive and for negative angles θ are averaged, the sine functions cancel, and only the cosine functions remain, while by subtracting the values of y_0 for positive and for negative angles, only the sine functions remain; that is,

$$\left. \begin{aligned} y_0(\theta) + y_0(-\theta) &= 2u(\theta), \\ y_0(\theta) - y_0(-\theta) &= 2v(\theta); \end{aligned} \right\} \quad . \quad . \quad . \quad . \quad . \quad (32)$$

hence, the cosine terms and the sine terms can be separated from each other by combining the instantaneous values of y_0 for positive angle θ and for negative angle $(-\theta)$, thus:

$$\left. \begin{aligned} u(\theta) &= \frac{1}{2} \{ y_0(\theta) + y_0(-\theta) \}, \\ v(\theta) &= \frac{1}{2} \{ y_0(\theta) - y_0(-\theta) \}. \end{aligned} \right\} \quad . \quad . \quad . \quad . \quad . \quad (33)$$

Usually, before separating the cosine and the sine terms, u and v , first the constant term a_0 is separated, as discussed above; that is, the alternating function $y_0 = y - a_0$ used. If the general periodic function y is used in equation (33), the constant term a_0 of this periodic function appears in the cosine term u , thus:

$$u(\theta) = \frac{1}{2}\{y(\theta) + y(-\theta)\} = a_0 + a_1 \cos \theta + a_2 \cos 2\theta + a_3 \cos 3\theta + \dots,$$

while $v(\theta)$ remains the same as when using y_0 .

86. Before separating the alternating function y_0 into the cosine function u and the sine function v , it usually is more convenient to resolve the alternating function y_0 into the odd series y_1 , and the even series y_2 , as discussed in the preceding paragraph, and then to separate y_1 and y_2 each into the cosine and the sine terms:

$$\left. \begin{aligned} u_1(\theta) &= \frac{1}{2}\{y_1(\theta) + y_1(-\theta)\} = a_1 \cos \theta + a_3 \cos 3\theta + a_5 \cos 5\theta + \dots; \\ v_1(\theta) &= \frac{1}{2}\{y_1(\theta) - y_1(-\theta)\} = b_1 \sin \theta + b_3 \sin 3\theta + b_5 \sin 5\theta + \dots; \end{aligned} \right\} \quad (34)$$

$$\left. \begin{aligned} u_2(\theta) &= \frac{1}{2}\{y_2(\theta) + y_2(-\theta)\} = a_2 \cos 2\theta + a_4 \cos 4\theta + \dots; \\ v_2(\theta) &= \frac{1}{2}\{y_2(\theta) - y_2(-\theta)\} = b_2 \sin 2\theta + b_4 \sin 4\theta + \dots \end{aligned} \right\} \quad (35)$$

In the odd functions u_1 and v_1 , a change from the negative angle $(-\theta)$ to the supplementary angle $(\pi - \theta)$ changes the angle of the trigonometric function by an odd multiple of π or 180 deg., that is, by a multiple of 2π or 360 deg., plus 180 deg., which signifies a reversal of the function, thus:

$$\left. \begin{aligned} u_1(\theta) &= \frac{1}{2}\{y_1(\theta) - y_1(\pi - \theta)\}, \\ v_1(\theta) &= \frac{1}{2}\{y_1(\theta) + y_1(\pi - \theta)\}. \end{aligned} \right\} \quad (36)$$

However, in the even functions u_2 and v_2 a change from the negative angle $(-\theta)$ to the supplementary angle $(\pi - \theta)$, changes the angles of the trigonometric function by an even multiple of π ; that is, by a multiple of 2π or 360 deg.; hence leaves the sign of the trigonometric function unchanged, thus:

$$\left. \begin{aligned} u_2(\theta) &= \frac{1}{2}\{y_2(\theta) + y_2(\pi - \theta)\}, \\ v_2(\theta) &= \frac{1}{2}\{y_2(\theta) - y_2(\pi - \theta)\}. \end{aligned} \right\} \quad (37)$$

To avoid the possibility of a mistake, it is preferable to use the relations (34) and (35), which are the same for the odd and for the even series.

87. Obviously, in the calculation of the constants a_n and b_n , instead of averaging from 0 to 180 deg., the average can be made from -90 deg. to $+90$ deg. In the cosine function $u(\theta)$, however, the same numerical values repeated with the same signs, from 0 to -90 deg., as from 0 to $+90$ deg., and the multipliers $\cos n\theta$ also have the same signs and the same numerical values from 0 to -90 deg., as from 0 to $+90$ deg. In the sine function, the same numerical values repeat from 0 to -90 deg., as from 0 to $+90$ deg., but with reversed signs, and the multipliers $\sin n\theta$ also have the same numerical values, but with reversed sign, from 0 to -90 deg., as from 0 to $+90$ deg. The products $u \cos n\theta$ and $v \sin n\theta$ thus traverse the same numerical values with the same signs, between 0 and -90 deg., as between 0 and $+90$ deg., and for deriving the averages, it thus is sufficient to average only from 0 to $\frac{\pi}{2}$, or 90 deg.; that is, over one quadrant.

Therefore, by resolving the periodic function y into the cosine components u and the sine components v , the calculation of the constants a_n and b_n is greatly simplified; that is, instead of averaging over one entire period, or 360 deg., it is necessary to average over only 90 deg., thus:

$$\left. \begin{aligned} a_1 &= 2 \text{ avg. } (u_1 \cos \theta)_0^{\frac{\pi}{2}}; & b_1 &= 2 \text{ avg. } (v_1 \sin \theta)_0^{\frac{\pi}{2}}; \\ a_2 &= 2 \text{ avg. } (u_2 \cos 2\theta)_0^{\frac{\pi}{2}}; & b_2 &= 2 \text{ avg. } (v_2 \sin 2\theta)_0^{\frac{\pi}{2}}; \\ a_3 &= 2 \text{ avg. } (u_3 \cos 3\theta)_0^{\frac{\pi}{2}}; & b_3 &= 2 \text{ avg. } (v_3 \sin 3\theta)_0^{\frac{\pi}{2}}; \\ a_4 &= 2 \text{ avg. } (u_4 \cos 4\theta)_0^{\frac{\pi}{2}}; & b_4 &= 2 \text{ avg. } (v_4 \sin 4\theta)_0^{\frac{\pi}{2}}; \\ a_5 &= 2 \text{ avg. } (u_5 \cos 5\theta)_0^{\frac{\pi}{2}}; & b_5 &= 2 \text{ avg. } (v_5 \sin 5\theta)_0^{\frac{\pi}{2}}; \\ &\text{etc.} & &\text{etc.} \end{aligned} \right\} \quad (38)$$

where u_1 is the cosine term of the odd function y_1 ; u_2 the cosine term of the even function y_2 ; u_3 is the cosine term of the odd function, after subtracting the term with $\cos \theta$; that is,

$$u_3 = u_1 - a_1 \cos \theta,$$

analogously, u_4 is the cosine term of the even function, after subtracting the term $\cos 2\theta$;

$$u_4 = u_2 - a_2 \cos 2\theta,$$

and in the same manner,

$$u_5 = u_3 - a_3 \cos 3\theta,$$

$$u_6 = u_4 - a_4 \cos 4\theta,$$

and so forth; v_1, v_2, v_3, v_4 , etc., are the corresponding sine terms.

When calculating the coefficients a_n and b_n by averaging over 90 deg., or over 180 deg. or 360 deg., it must be kept in mind that the terminal values of y respectively of u or v , that is, the values for $\theta=0$ and $\theta=90$ deg. (or $\theta=180$ deg. or 360 deg. respectively) are to be taken as one-half only, since they are the ends of the measured area of the curves $a_n \cos n\theta$ and $b_n \sin n\theta$, which area gives as twice its average height the values a_n and b_n , as discussed in the preceding.

In resolving an empirical periodic function into a trigonometric series, just as in most engineering calculations, the most important part is to arrange the work so as to derive the results expeditiously and rapidly, and at the same time accurately. By proceeding, for instance, immediately by the general method, equations (17) and (18), the work becomes so extensive as to be a serious waste of time, while by the systematic resolution into simpler functions the work can be greatly reduced.

88. In resolving a general periodic function $y(\theta)$ into a trigonometric series, the most convenient arrangement is:

1. To separate the constant term a_0 , by averaging all the instantaneous values of $y(\theta)$ from 0 to 360 deg. (counting the end values at $\theta=0$ and at $\theta=360$ deg. one half, as discussed above):

$$a_0 = \text{avg. } \{y(\theta)\}_{0^{2\pi}}, \quad . \quad . \quad . \quad . \quad . \quad (10)$$

and then subtracting a_0 from $y(\theta)$, gives the alternating function,

$$y_0(\theta) = y(\theta) - a_0.$$

2. To resolve the general alternating function $y_0(\theta)$ into the odd function $y_1(\theta)$, and the even function $y_2(\theta)$,

$$y_1(\theta) = \frac{1}{2} \{ y_0(\theta) - y_0(\theta + \pi) \}; \quad . \quad . \quad . \quad . \quad (23)$$

$$y_2(\theta) = \frac{1}{2} \{ y_0(\theta) + y_0(\theta + \pi) \}. \quad . \quad . \quad . \quad . \quad (22)$$

3. To resolve $y_1(\theta)$ and $y_2(\theta)$ into the cosine terms u and the sine terms v ,

$$\left. \begin{aligned} u_1(\theta) &= \frac{1}{2} \{ y_1(\theta) + y_1(-\theta) \}; \\ v_1(\theta) &= \frac{1}{2} \{ y_1(\theta) - y_1(-\theta) \}; \end{aligned} \right\} . \quad . \quad . \quad . \quad (34)$$

$$\left. \begin{aligned} u_2(\theta) &= \frac{1}{2} \{ y_2(\theta) + y_2(-\theta) \}; \\ v_2(\theta) &= \frac{1}{2} \{ y_2(\theta) - y_2(-\theta) \}. \end{aligned} \right\} . \quad . \quad . \quad . \quad (35)$$

4. To calculate the constants $a_1, a_2, a_3, \dots; b_1, b_2, b_3, \dots$ by the averages,

$$\left. \begin{aligned} a_n &= 2 \text{ avg. } (u_n \cos n\theta)_0^{\frac{\pi}{2}}; \\ b_n &= 2 \text{ avg. } (v_n \sin n\theta)_0^{\frac{\pi}{2}}. \end{aligned} \right\} . \quad . \quad . \quad . \quad (38)$$

If the periodic function is known to contain no even harmonics, that is, is a symmetrical alternating wave, steps 1 and 2 are omitted.

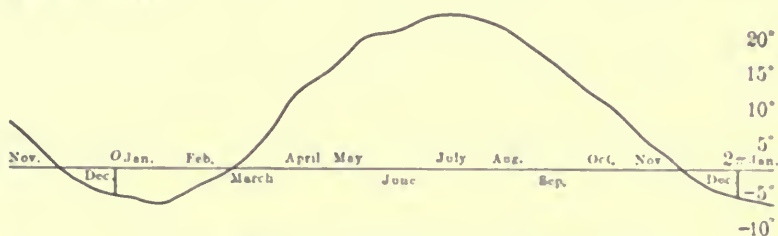


FIG. 45. Mean Daily Temperature at Schenectady.

89. As illustration of the resolution of a general periodic wave may be shown the resolution of the observed mean daily temperatures of Schenectady throughout the year, as shown in Fig. 45, up to the 7th harmonic.*

* The numerical values of temperature cannot claim any great absolute accuracy, as they are averaged over a relatively small number of years only, and observed by instruments of only moderate accuracy. For the purpose of illustrating the resolution of the empirical curve into a trigonometric series, this is not essential, however.

TABLE IV

(1) θ	(2) y	(3) $y - a_0 = y_0$	(4) y_1	(5) y_2
Jan. 0	- 4.2	-12.95	-13.10	+0.15
10	- 4.7	-13.45	-13.55	+0.10
20	- 5.2	-13.95	-13.65	-0.30
Feb. 30	- 5.4	-14.15	-13.55	-0.60
40	- 3.8	-12.55	-12.35	-0.20
50	- 2.6	-11.35	-11.20	-0.15
Mar. 60	- 1.6	-10.35	- 9.75	-0.60
70	+ 0.2	- 8.55	- 7.65	-0.90
80	+ 1.8	- 6.95	- 6.05	-0.90
Apr. 90	+ 5.1	- 3.65	- 3.35	-0.30
100	+ 9.1	+ 0.35	- 0.35	+0.70
110	+11.5	+ 2.75	+ 1.75	+1.00
May 120	+13.3	+ 4.55	+ 3.90	+0.65
130	+15.2	+ 6.45	+ 5.85	+0.60
140	+17.7	+ 8.95	+ 8.15	+0.80
June 150	+19.2	+10.45	+10.10	+0.35
160	+19.5	+10.75	+10.80	-0.05
170	+20.6	+11.85	+12.15	-0.30
July 180	+22.0	+13.25		
190	+22.4	+13.65		
200	+22.1	+13.35		
Aug. 210	+21.7	+12.95		
220	+20.9	+12.15		
230	+19.8	+11.05		
Sept. 240	+17.9	+ 9.15		
250	+15.5	+ 6.75		
260	+13.8	+ 5.15		
Oct. 270	+11.8	+ 3.05		
280	+ 9.8	+ 1.05		
290	+ 8.0	- 0.75		
Nov. 300	+ 5.5	- 3.25		
310	+ 3.5	- 5.25		
320	+ 1.4	- 7.35		
Dec. 330	- 1.0	- 9.75		
340	- 2.1	-10.85		
350	- 3.7	-12.45		
Total.....	315.1			
Divided by 36.	8.75 = a_0			

TABLE V.

(1) θ	(2) y_1	(3) u_1	(4) v_1	(5) y_2	(6) u_2	(7) v_2
-90	+ 3.35			-0.30		
-80	+ 0.35			+0.70		
-70	- 1.75			+1.00		
-60	- 3.90			+0.65		
-50	- 5.85			+0.60		
-40	- 8.15			+0.80		
-30	-10.10			+0.35		
-20	-10.80			-0.05		
-10	-12.15			-0.30		
0	-13.10	-13.10	0	+0.15	+0.15	0
+10	-13.55	-12.85	-0.70	+0.10	-0.10	+0.20
+20	-13.65	-12.23	-1.42	-0.30	-0.17	-0.12
+30	-13.55	-11.82	-1.73	-0.60	-0.12	-0.47
+40	-12.35	-10.25	-2.10	-0.20	+0.30	-0.50
+50	-11.20	- 8.53	-2.67	-0.15	+0.22	-0.37
+60	- 9.75	- 6.82	-2.93	-0.60	+0.02	-0.62
+70	- 7.65	- 4.70	-2.95	-0.90	+0.05	-0.95
+80	- 6.05	- 2.85	-3.20	-0.90	-0.10	-0.80
+90	- 3.35	0	-3.35	-0.30	-0.30	0

TABLE VI.
COSINE SERIES u_i .

(1) θ	(2) u_1	(3) $\cos \theta$	(4) $u_1 \cos \theta$	(5) $a_1 \cos \theta$	(6) u_3	(7) $u_3 \cos 3\theta$	(8) $a_3 \cos 3\theta$	(9) u_5	(10) $u_5 \cos 5\theta$	(11) $u_5 \cos 7\theta$
0	-13.10	1	-13.10($\times \frac{1}{2}$)	-13.28	+0.18	+0.18($\times \frac{1}{2}$)	+0.33	-0.15	-0.15($\times \frac{1}{2}$)	-0.15($\times \frac{1}{2}$)
10	-12.85	0.985	-12.65	-13.05	-0.20	+0.17	0.285	-0.085	-0.054	-0.029
20	-12.23	0.940	-11.50	-12.48	+0.25	+0.12	0.165	+0.085	-0.015	-0.065
30	-11.82	0.866	-10.25	-11.50	-0.32	0	0	-0.32	+0.277	-0.239
40	-10.25	0.766	-7.83	-10.15	-0.10	+0.05	-0.165	+0.065	-0.061	-0.011
50	-8.53	0.643	-5.46	-8.00	-0.53	+0.46	-0.285	-0.245	+0.084	+0.083
60	-6.82	0.5	-3.41	-6.64	-0.18	+0.18	-0.33	+0.15	+0.075	+0.037
70	-4.70	0.342	-1.61	-4.54	-0.16	+0.14	-0.285	+0.125	+0.123	-0.079
80	-2.85	0.174	-0.50	-2.30	-0.55	+0.27	-0.165	-0.385	-0.293	+0.276
90	0	0	0	0	0	0	0	0	0	0
Total			-59.75						+0.061	-0.101
Divided by 9			-6.64						+0.0068	-0.011
Multipled by 2			-13.28= a_1						+0.014= a_5	-0.022= a_7

TABLE VII.
SINE SERIES v_1 .

(1) θ	(2) v_1	(3) $\sin \theta$	(4) $v_1 \sin \theta$	(5) $b_1 \sin \theta$	(6) v_1	(7) $v_1 \sin 3\theta$	(8) $b_3 \sin 3\theta$	(9) v_1	(10) $v_1 \sin 5\theta$	(11) $v_1 \sin 7\theta$
0	0	0	0	0	0	0	0	0	0	0
10	-0.70	0.174	-0.122	-0.38	-0.32	-0.16	-0.07	-0.25	-0.19	-0.23
20	-1.42	0.342	-0.485	-1.14	-0.28	-0.24	-0.12	-0.16	-0.16	-0.10
30	-1.73	0.5	-0.865	-1.66	-0.07	-0.67	-0.14	+0.07	+0.035	-0.035
40	-2.10	0.643	-1.345	-2.13	+0.03	+0.03	-0.12	0.15	-0.05	-0.15
50	-2.67	0.766	-2.040	-2.55	-0.12	-0.06	-0.07	-0.05	+0.05	+0.01
60	-2.93	0.866	-2.530	-2.88	-0.05	0	0	-0.05	+0.04	-0.04
70	-2.95	0.940	-2.770	-3.13	+0.18	-0.09	+0.07	+0.11	-0.02	+0.085
80	-3.24	0.985	-3.150	-3.28	+0.08	-0.07	+0.12	-0.04	-0.025	+0.01
90	3.35	1	-3.35($\times 4$)	-3.33	-0.02	+0.02($\times 4$)	+0.14	-0.16	-0.16($\times 4$)	+0.16($\times 4$)
Total ...			14.982			-0.65			-0.40	-0.37
Divided by 9 ...			-1.665			-0.07			-0.044	-0.041
Multiplied by 2 ...			-3.33 = b_1			-0.14 = b_3			-0.09 = b_5	-0.082 = b_7

TABLE VIII.
COSINE SERIES u_2 .

(1) θ	(2) u_2	(3) $u_2 \cos 2\theta$	(4) $a_2 \cos 2\theta$	(5) u_4	(6) $u_4 \cos 4\theta$	(7) $a_4 \cos 4\theta$	(8) u_6	$u_6 \cos 6\theta$
0	+0.15	$\frac{1}{2}(+0.15)$	0	+0.15	$\frac{1}{2}(+0.15)$	-0.16	+0.31	$\frac{1}{2}(+0.31)$
10	-0.10	-0.09		-0.10	-0.08	-0.12	+0.02	+0.01
20	-0.17	-0.13		-0.17	-0.03	-0.03	-0.14	+0.07
30	-0.12	-0.06		-0.12	+0.06	+0.08	-0.20	+0.20
40	+0.30	+0.05		+0.30	-0.29	+0.15	+0.15	-0.07
50	+0.22	-0.04		+0.22	-0.21	+0.15	+0.07	+0.03
60	+0.02	-0.01		+0.02	-0.01	+0.08	-0.06	-0.06
70	+0.05	-0.04		+0.05	+0.01	-0.03	+0.08	+0.04
80	-0.10	+0.09		-0.10	-0.08	-0.12	+0.02	-0.01
90	-0.30	$\frac{1}{2}(+0.30)$	0	-0.30	$\frac{1}{2}(+0.30)$	-0.16	-0.14	$\frac{1}{2}(+0.14)$
Total.....		-0.01			-0.71			+0.44
Divided by 9.....		-0.001			-0.079			+0.049
Multiplied by 2....		-0.002			-0.158			+0.098
		$=a_2$			$=a_3$			$=a_6$

TABLE IX.
SINE SERIES v_2 .

(1) θ	(2) v_2	(3) $v_2 \sin 2\theta$	(4) $b_2 \sin 2\theta$	(5) v_4	(6) $v_4 \sin 4\theta$	(7) $b_4 \sin 4\theta$	(8) v_6	(9) $v_6 \sin 6\theta$
0	0	0						
10	+0.20	+0.07	-0.20	+0.40	+0.26	+0.22	+0.18	+0.16
20	-0.12	-0.08	-0.39	+0.27	+0.27	+0.34	-0.07	-0.07
30	-0.47	-0.41	-0.52	+0.05	+0.04	+0.30	-0.25	+0
40	-0.50	-0.49	-0.59	+0.09	+0.03	+0.12	-0.03	+0.03
50	-0.37	-0.36	-0.59	+0.22	-0.08	-0.12	+0.34	-0.30
60	-0.62	-0.54	-0.52	-0.10	+0.09	-0.30	+0.20	0
70	-0.95	-0.61	-0.39	-0.56	+0.55	-0.34	-0.22	-0.19
80	-0.80	-0.27	-0.20	-0.60	+0.39	+0.22	-0.38	-0.33
90	0	0						
Total.....		-2.69			+1.55			-0.70
Divided by 9		-0.30			+0.172			-0.078
Divided by 2		-0.60			+0.344			-0.156
		$=b_2$			$=b_4$			$=b_6$

Table IV gives the resolution of the periodic temperature function into the constant term a_0 , the odd series y_1 and the even series y_2 .

Table V gives the resolution of the series y_1 and y_2 into the cosine and sine series u_1, v_1, u_2, v_2 .

Tables VI to IX give the resolutions of the series u_1, v_1, u_2, v_2 , and thereby the calculation of the constants a_n and b_n .

90. The resolution of the temperature wave, up to the 7th harmonic, thus gives the coefficients:

$$\begin{aligned} a_0 &= +8.75; \\ a_1 &= -13.28; & b_1 &= -3.33; \\ a_2 &= -0.001; & b_2 &= -0.602; \\ a_3 &= -0.33; & b_3 &= -0.14; \\ a_4 &= -0.154; & b_4 &= +0.386; \\ a_5 &= +0.014; & b_5 &= -0.090; \\ a_6 &= +0.100; & b_6 &= -0.154; \\ a_7 &= -0.022; & b_7 &= -0.082; \end{aligned}$$

or, transforming by the binomial, $a_n \cos n\theta + b_n \sin n\theta = c_n \cos(n\theta - \gamma_n)$, by substituting $c_n = \sqrt{a_n^2 + b_n^2}$ and $\tan \gamma_n = \frac{b_n}{a_n}$ gives,

$$\begin{aligned} c_0 &= +8.75; \\ c_1 &= -13.69; \quad \gamma_1 = +14.15^\circ; \text{ or } \gamma_1 = +14.15^\circ; \\ c_2 &= -0.602; \quad \gamma_2 = +89.9^\circ; \text{ or } \frac{\gamma_2}{2} = +44.95^\circ + 180n; \\ c_3 &= +0.359; \quad \gamma_3 = -23.0^\circ; \text{ or } \frac{\gamma_3}{3} = -7.7 + 120n = +112.3 + 120m; \\ c_4 &= -0.416; \quad \gamma_4 = -68.2^\circ; \text{ or } \frac{\gamma_4}{4} = -17.05 + 90n = +72.95 + 90m; \\ c_5 &= +0.091; \quad \gamma_5 = -81.15^\circ; \text{ or } \frac{\gamma_5}{5} = -16.23 + 72n = +55.77 + 72m; \\ c_6 &= +0.184; \quad \gamma_6 = -57.0^\circ; \text{ or } \frac{\gamma_6}{6} = -9.5 + 60n = +50.5 + 60m; \\ c_7 &= -0.085; \quad \gamma_7 = +75.0^\circ; \text{ or } \frac{\gamma_7}{7} = +10.7 + 51.4n, \end{aligned}$$

where n and m may be any integer number.

Since to an angle γ_n , any multiple of 2π or 360 deg. may be added, any multiple of $\frac{360}{n}$ may be added to the angle $\frac{\gamma_n}{n}$, and thus the angle $\frac{\gamma_n}{n}$ may be made positive, etc.

91. The equation of the temperature wave thus becomes:

$$\begin{aligned} y = & 8.75 - 13.69 \cos (\theta - 14.15^\circ) - 0.602 \cos 2(\theta - 44.95^\circ) \\ & - 0.359 \cos 3(\theta - 52.3^\circ) - 0.416 \cos 4(\theta - 72.95^\circ) \\ & - 0.091 \cos 5(\theta - 19.77^\circ) - 0.184 \cos 6(\theta - 20.5^\circ) \\ & - 0.085 \cos 7(\theta - 10.7^\circ); \end{aligned} \quad (a)$$

or, transformed to sine functions by the substitution,

$$\cos \omega = -\sin (\omega - 90^\circ):$$

$$\begin{aligned} y = & 8.75 + 13.69 \sin (\theta - 104.15^\circ) + 0.602 \sin 2(\theta - 89.95^\circ) \\ & + 0.359 \sin 3(\theta - 82.3^\circ) + 0.416 \sin 4(\theta - 95.45^\circ) \\ & + 0.091 \sin 5(\theta - 109.77^\circ) + 0.184 \sin 6(\theta - 95.5^\circ) \\ & + 0.085 \sin 7(\theta - 75^\circ). \end{aligned} \quad (b)$$

The cosine form is more convenient for some purposes, the sine form for other purposes.

Substituting $\beta = \theta - 14.15^\circ$; or, $\delta = \theta - 104.15^\circ$, these two equations (a) and (b) can be transformed into the form,

$$\begin{aligned} y = & 8.75 - 13.69 \cos \beta - 0.62 \cos 2(\beta - 30.8^\circ) - 0.359 \cos 3(\beta - 38.15^\circ) \\ & - 0.416 \cos 4(\beta - 58.8^\circ) - 0.091 \cos 5(\beta - 5.6^\circ) \\ & - 0.184 \cos 6(\beta - 6.35^\circ) - 0.085 \cos 7(\beta - 48.0^\circ), \end{aligned} \quad (c)$$

and

$$\begin{aligned} y = & 8.75 + 13.69 \sin \delta + 0.602 \sin 2(\delta + 14.2^\circ) + 0.359 \sin 3(\delta + 21.85^\circ) \\ & + 0.416 \sin 4(\delta + 8.7^\circ) + 0.91 \sin 5(\delta - 5.6^\circ) \\ & + 0.184 \sin 6(\delta + 8.65^\circ) + 0.085 \sin 7(\delta + 29.15^\circ). \end{aligned} \quad (d)$$

The periodic variation of the temperature y , as expressed by these equations, is a result of the periodic variation of the thermomotive force; that is, the solar radiation. This latter

is a minimum on Dec. 22d, that is, 9 time-degrees before the zero of θ , hence may be expressed approximately by:

$$z = c - h \cos (\theta + 9^\circ);$$

or substituting β respectively δ for θ :

$$z = c - h \cos (\beta + 23.15^\circ)$$

$$= c + h \sin (\delta + 23.15^\circ).$$

This means: the maximum of y occurs 23.15 deg. after the maximum of z ; in other words, the temperature lags 23.15 deg., or about $\frac{1}{4}$ period, behind the thermomotive force.

Near $\delta = 0$, all the sine functions in (d) are increasing; that is, the temperature wave rises steeply in spring.

Near $\delta = 180$ deg., the sine functions of the odd angles are decreasing, of the even angles increasing, and the decrease of the temperature wave in fall thus is smaller than the increase in spring.

The fundamental wave greatly preponderates, with amplitude $c_1 = 13.69$.

In spring, for $\delta = -14.5$ deg., all the higher harmonics rise in the same direction, and give the sum 1.74, or 12.7 per cent of the fundamental. In fall, for $\delta = -14.5 + \pi$, the even harmonics decrease, the odd harmonics increase the steepness, and give the sum -0.67 , or -4.9 per cent.

Therefore, in spring, the temperature rises 12.7 per cent faster, and in autumn it falls 4.9 per cent slower than corresponds to a sine wave, and the difference in the rate of temperature rise in spring, and temperature fall in autumn thus is $12.7 + 4.9 = 17.6$ per cent.

The maximum rate of temperature rise is $90 - 14.5 = 75.5$ deg. behind the temperature minimum, and $23.15 + 75.5 = 98.7$ deg. behind the minimum of the thermomotive force.

As most periodic functions met by the electrical engineer are symmetrical alternating functions, that is, contain only the odd harmonics, in general the work of resolution into a trigonometric series is very much less than in above example. Where such reduction has to be carried out frequently, it is advisable to memorize the trigonometric functions, from 10 to 10 deg., up to 3 decimals: that is, within the accuracy of the slide rule, as thereby the necessity of looking up tables is

eliminated and the work therefore done much more expeditiously. In general, the slide rule can be used for the calculations.

As an example of the simpler reduction of a symmetrical alternating wave, the reader may resolve into its harmonics, up to the 7th, the exciting current of the transformer, of which the numerical values are given, from 10 to 10 deg. in Table X.

C. REDUCTION OF TRIGONOMETRIC SERIES BY POLYPHASE RELATION.

92. In some cases the reduction of a general periodic function, as a complex wave, into harmonics can be carried out in a much quicker manner by the use of the polyphase equation, Chapter III, Part A (23). Especially is this true if the complete equation of the trigonometric series, which represents the periodic function, is not required, but the existence and the amount of certain harmonics are to be determined, as for instance whether the periodic function contain even harmonics or third harmonics, and how large they may be.

This method does not give the coefficients a_n , b_n of the individual harmonics, but derives from the numerical values of the general wave the numerical values of any desired harmonic. This harmonic, however, is given together with all its multiples; that is, when separating the third harmonic, in it appears also the 6th, 9th, 12th, etc.

In separating the even harmonics y_2 from the general wave y , in paragraph 84, by taking the average of the values of y for angle θ , and the values of y for angles $(\theta + \pi)$, this method has already been used.

Assume that to an angle θ there is successively added a constant quantity a , thus: θ ; $\theta + a$; $\theta + 2a$; $\theta + 3a$; $\theta + 4a$, etc., until the same angle θ plus a multiple of 2π is reached;

$\theta + na = \theta + 2m\pi$; that is, $a = \frac{2m\pi}{n}$; or, in other words, a is

$1/n$ of a multiple of 2π . Then the sum of the cosine as well as the sine functions of all these angles is zero:

$$\begin{aligned} \cos \theta + \cos (\theta + a) + \cos (\theta + 2a) + \cos (\theta + 3a) + \dots \\ + \cos (\theta + [n-1]a) = 0; \end{aligned} \quad (1)$$

$$\begin{aligned} \sin \theta + \sin (\theta + a) + \sin (\theta + 2a) + \sin (\theta + 3a) + \dots \\ + \sin (\theta + [n-1]a) = 0, \quad (2) \end{aligned}$$

where

$$na = 2m\pi. \quad (3)$$

These equations (1) and (2) hold for all values of a , except for $a = 2\pi$, or a multiple thereof. For $a = 2\pi$ obviously all the terms of equation (1) or (2) become equal, and the sums become $n \cos \theta$ respectively $n \sin \theta$.

Thus, if the series of numerical values of y is divided into n successive sections, each covering $\frac{2\pi}{n}$ degrees, and these sections added together,

$$\begin{aligned} y(\theta) + y\left(\theta + \frac{2\pi}{n}\right) + y\left(\theta + 2\frac{2\pi}{n}\right) + y\left(\theta + 3\frac{2\pi}{n}\right) + \dots \\ + y\left(\theta + [n-1]\frac{2\pi}{n}\right), \quad (4) \end{aligned}$$

In this sum, all the harmonics of the wave y cancel by equations (1) and (2), except the n th harmonic and its multiples,

$$a_n \cos n\theta + b_n \sin n\theta; \quad a_{2n} \cos 2n\theta + b_{2n} \sin 2n\theta, \text{ etc.}$$

in the latter all the terms of the sum (4) are equal; that is, the sum (4) equals n times the n th harmonic, and its multiples. Therefore, the n th harmonic of the periodic function y , together with its multiples, is given by

$$y_n(\theta) = \frac{1}{n} \left\{ y(\theta) + y\left(\theta + \frac{2\pi}{n}\right) + y\left(\theta + 2\frac{2\pi}{n}\right) + \dots + y\left(\theta + [n-1]\frac{2\pi}{n}\right) \right\} \quad (5)$$

For instance, for $n = 2$,

$$y_2 = \frac{1}{2} \{ y(\theta) + y(\theta + \pi) \},$$

gives the sum of all the even harmonics; that is, gives the second harmonic together with its multiples, the 4th, 6th, etc., as seen in paragraph 7, and for, $n = 3$,

$$y_3 = \frac{1}{3} \left\{ y(\theta) + y\left(\theta + \frac{2\pi}{3}\right) + y\left(\theta + \frac{4\pi}{3}\right) \right\},$$

gives the third harmonic, together with its multiples, the 6th, 9th, etc.

This method does not give the mathematical expression of the harmonics, but their numerical values. Thus, if the mathematical expressions are required, each of the component harmonics has to be reduced from its numerical values to the mathematical equation, and the method then usually offers no advantage.

It is especially suitable, however, where certain classes of harmonics are desired, as the third together with its multiples. In this case from the numerical values the effective value, that is, the equivalent sine wave may be calculated.

93. As illustration may be investigated the separation of the third harmonics from the exciting current of a transformer.

TABLE X

A						
(1) θ	(2) i	(3) θ	(4) i	(5) θ	(6) i	(7) i_3
0	+24.0	120	-15.1	240	+8.5	+5.8
10	+20.0	130	-16.5	250	+10	+4.5
20	+12	140	-18.5	260	+11	+1.5
30	+4	150	-21	270	+12	-1.7
40	-1.5	160	-22.7	280	+13	-3.7
50	-6.5	170	-23.7	290	+14	-5.4
60	-8.5	180	-24	300	+15.1	-5.8
B						
3θ	i_3	3θ	i_3	3θ	i_3	i_9
0	+5.8	120	-3.7	240	-1.5	+0.2
30	+4.5	150	-5.4	270	+1.7	+0.3
60	+1.5	180	-5.8	300	+3.7	-0.2

'In table X A, are given, in columns 1, 3, 5, the angles θ , from 10 deg. to 10 deg., and in columns 2, 4, 6, the corresponding values of the exciting current i , as derived by calculation from the hysteresis cycle of the iron, or by measuring from the

photographic film of the oscillograph. Column 7 then gives one-third the sum of columns 2, 4, and 6, that is, the third harmonic with its overtones, i_3 .

To find the 9th harmonic and its overtones i_9 , the same method is now applied to i_3 , for angle 3θ . This is recorded in Table X B.

In Fig. 46 are plotted the total exciting current i , its third harmonic i_3 , and the 9th harmonic i_9 .

This method has the advantage of showing the limitation of the exactness of the results resulting from the limited num-

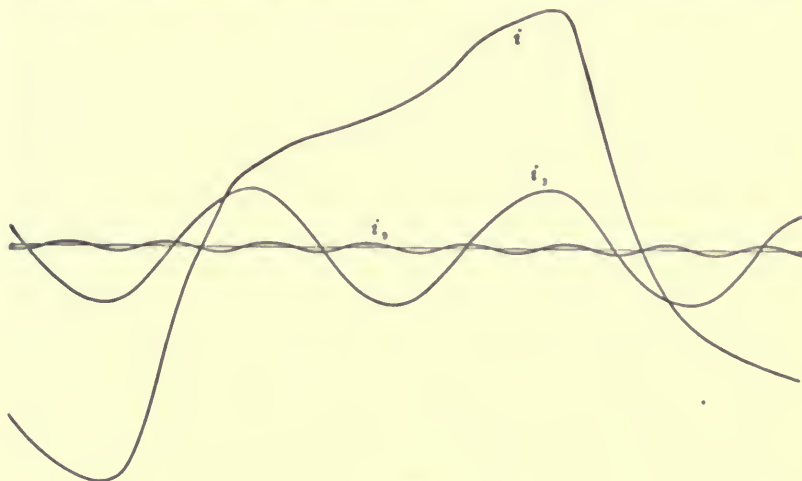


FIG. 46.

ber of numerical values of i , on which the calculation is based. Thus, in the example, Table X, in which the values of i are given for every 10 deg., values of the third harmonic are derived for every 30 deg., and for the 9th harmonic for every 90 deg.; that is, for the latter, only two points per half wave are determinable from the numerical data, and as the two points per half wave are just sufficient to locate a sine wave, it follows that within the accuracy of the given numerical values of i , the 9th harmonic is a sine wave, or in other words, to determine whether still higher harmonics than the 9th exist, requires for i more numerical values than for every 10 deg.

As further practice, the reader may separate from the gen-

eral wave of current, i_0 in Table XI, the even harmonics i_2 , by above method,

$$i_2 = \frac{1}{2}\{i_0(\theta) + i_0(\theta + 180 \text{ deg.})\},$$

and also the sum of the odd harmonics, as the residue,

$$i_1 = i_0 - i_2,$$

then from the odd harmonics i_1 may be separated the third harmonic and its multiples,

$$i_3 = \frac{1}{3}\{i_1(\theta) + i_1(\theta + 120 \text{ deg.}) + i_1(\theta + 240 \text{ deg.})\},$$

and in the same manner from i_3 may be separated its third harmonic; that is, i_9 .

Furthermore, in the sum of even harmonics, from i_2 may again be separated its second harmonic, i_4 , and its multiples, and therefrom, i_8 , and its third harmonic, i_6 , and its multiples, thus giving all the harmonics up to the 9th, with the exception of the 5th and the 7th. These latter two would require plotting the curve and taking numerical values at different intervals, so as to have a number of numerical values divisible by 5 or 7.

It is further recommended to resolve this unsymmetrical exciting current of Table XI into the trigonometric series by calculating the coefficients a_n and b_n , up to the 7th, in the manner discussed in paragraphs 6 to 8.

TABLE XI

θ	i_0	θ	i_0	θ	i_0	θ	i_0
0	+95.7	90	-26.7	180	-34.3	270	- 3.3
10	+78.7	100	-27.3	190	-27.3	280	- 1.8
20	+53.7	110	-28.1	200	-16.8	290	+ 1.2
30	+23.7	120	-28.8	210	-11.3	300	+ 4.7
40	- 2.3	130	-29.3	220	- 8.3	310	+10.7
50	-16.3	140	-29.8	230	- 7.3	320	+22.7
60	-22.8	150	-31	240	- 6.3	330	+41.7
70	-24.3	160	-32.6	250	- 5.3	340	+65.7
80	-25.8	170	-33.8	260	- 4.3	350	+85.7

D. CALCULATION OF TRIGONOMETRIC SERIES FROM OTHER TRIGONOMETRIC SERIES.

94. An hydraulic generating station has for a long time been supplying electric energy over moderate distances, from a number of 750-kw. 4400-volt 60-cycle three-phase generators. The station is to be increased in size by the installation of some larger modern three-phase generators, and from this station 6000 kw. are to be transmitted over a long distance transmission line at 44,000 volts. The transmission line has a length of 60 miles, and consists of three wires No. 0 B. & S. with 5 ft. between the wires.

The question arises, whether during times of light load the old 750-kw. generators can be used economically on the transmission line. These old machines give an electromotive force wave, which, like that of most earlier machines, differs considerably from a sine wave, and it is to be investigated, whether, due to this wave-shape distortion, the charging current of the transmission line will be so greatly increased over the value which it would have with a sine wave of voltage, as to make the use of these machines on the transmission line uneconomical or even unsafe.

Oscillograms of these machines, resolved into a trigonometric series, give for the voltage between each terminal and the neutral, or the Y voltage of the three-phase system, the equation:

$$e = e_0 \{ \sin \theta - 0.12 \sin (3\theta - 2.3^\circ) - 0.23 \sin (5\theta - 1.5^\circ) + 0.13 \sin (7\theta - 6.2^\circ) \}. \quad (1)$$

In first approximation, the line capacity may be considered as a condenser shunted across the middle of the line; that is, half the line resistance and half the line reactance is in series with the line capacity.

As the receiving apparatus do not utilize the higher harmonics of the generator wave, when using the old generators, their voltage has to be transformed up so as to give the first harmonic or fundamental of 44,000 volts.

44,000 volts between the lines (or delta) gives 44,000 : $\sqrt{3}$ 25,400 volts between line and neutral. This is the effective

value, and the maximum value of the fundamental voltage wave thus is: $25,400 \times \sqrt{2} = 36,000$ volts, or 36 kv.; that is, $e_0 = 36$, and

$$e = 36 \{ \sin \theta - 0.12 \sin (3\theta - 2.3^\circ) - 0.23 \sin (5\theta - 1.5^\circ) + 0.13 \sin (7\theta - 6.2^\circ) \}, \quad (2)$$

would be the voltage supplied to the transmission line at the high potential terminals of the step-up transformers.

From the wire tables, the resistance per mile of No. 0 B. & S. copper line wire is $r_0 = 0.52$ ohm.

The inductance per mile of wire is given by the formula:

$$L_0 = 0.7415 \log \frac{l_s}{l_r} + 0.0805 \text{mh}, \quad (3)$$

where l_s is the distance between the wires, and l_r the radius of the wire.

In the present case, this gives $l_s = 5 \text{ ft.} = 60 \text{ in.}$ $l_r = 0.1625 \text{ in.}$ $L_0 = 1.9655 \text{ mh.}$, and, herefrom it follows that the reactance, at $f = 60$ cycles is

$$x_0 = 2\pi f L_0 = 0.75 \text{ ohms per mile.} \quad (4)$$

The capacity per mile of wire is given by the formula:

$$C_0 = \frac{0.0408}{\log \frac{l_s}{l_r}} \text{ mf.;} \quad (5)$$

hence, in the present case, $C_0 = 0.0159 \text{ mf.}$, and the condensive reactance is derived herefrom as:

$$x_{c0} = \frac{1}{2\pi f C_0} = 166000 \text{ ohms;} \quad (6)$$

60 miles of line then give the condensive reactance,

$$x_c = \frac{x_{c0}}{60} = 2770 \text{ ohms;} \quad (7)$$

30 miles, or half the line (from the generating station to the middle of the line, where the line capacity is represented by a shunted condenser) give: the resistance, $r = 30r_0 = 15.6$ ohms;

the inductive reactance, $x=30x_0=22.5$ ohms, and the equivalent circuit of the line now consists of the resistance r , inductive reactance x and condensive reactance x_c , in series with each other in the circuit of the supply voltage e .

95. If i = current in the line (charging current) the voltage consumed by the line resistance r is ri .

The voltage consumed by the inductive reactance x is $x \frac{di}{d\theta}$;

the voltage consumed by the condensive reactance x_c is $x_c \int i d\theta$, and therefore,

$$e = x \frac{di}{d\theta} + ri + x_c \int i d\theta. \quad (7)$$

Differentiating this equation, for the purpose of eliminating the integral, gives

$$\left. \begin{aligned} \frac{de}{d\theta} &= x \frac{d^2i}{d\theta^2} + r \frac{di}{d\theta} + x_c i; \\ \text{or} \quad \frac{de}{d\theta} &= 22.5 \frac{d^2i}{d\theta^2} + 16.6 \frac{di}{d\theta} + 2770i. \end{aligned} \right\} \quad (8)$$

The voltage e is given by (2), which equation, by resolving the trigonometric functions, gives

$$e = 36 \sin \theta - 4.32 \sin 3\theta - 8.28 \sin 5\theta + 1.64 \sin 7\theta \\ + 0.18 \cos 3\theta + 0.22 \cos 5\theta - 0.50 \cos 7\theta; \quad (9)$$

hence, differentiating,

$$\frac{de}{d\theta} = 36 \cos \theta - 12.96 \cos 3\theta - 41.4 \cos 5\theta + 32.5 \cos 7\theta \\ - 0.54 \sin 3\theta - 1.1 \sin 5\theta + 3.5 \sin 7\theta, \quad (10)$$

Assuming now for the current i a trigonometric series with indeterminate coefficients,

$$i = a_1 \cos \theta + a_3 \cos 3\theta + a_5 \cos 5\theta + a_7 \cos 7\theta \\ + b_1 \sin \theta + b_3 \sin 3\theta + b_5 \sin 5\theta + b_7 \sin 7\theta, \quad (11)$$

substitution of (10) and (11) into equation (8) must give an identity, from which equations for the determination of a_n and b_n are derived; that is, since the product of substitution must be an identity, all the factors of $\cos \theta$, $\sin \theta$, $\cos 3\theta$, $\sin 3\theta$, etc., must vanish, and this gives the eight equations:

$$\left. \begin{aligned} 36 &= 2770a_1 + 15.6b_1 - 22.5a_1; \\ 0 &= 2770b_1 - 15.6a_1 - 22.5b_1; \\ -12.96 &= 2770a_3 + 46.8b_3 - 202.5a_3; \\ -0.54 &= 2770b_3 - 46.8a_3 - 202.5b_3; \\ -41.4 &= 2770a_5 + 78b_5 - 562.5a_5; \\ -1.1 &= 2770b_5 - 78a_5 - 56.25b_5; \\ 32.5 &= 2770a_7 + 109.2b_7 - 1102.5a_7; \\ 3.5 &= 2770b_7 - 109.2a_7 - 1102.5b_7. \end{aligned} \right\} \quad (12)$$

Resolved, these equations give

$$\left. \begin{aligned} a_1 &= 13.12; \\ b_1 &= 0.07; \\ a_3 &= -5.03; \\ b_3 &= -0.30; \\ a_5 &= -18.72; \\ b_5 &= -1.15; \\ a_7 &= 19.30; \\ b_7 &= 3.37; \end{aligned} \right\} \quad (13)$$

hence,

$$\left. \begin{aligned} i &= 13.12 \cos \theta - 5.03 \cos 3\theta - 18.72 \cos 5\theta + 19.30 \cos 7\theta \\ &+ 0.07 \sin \theta - 0.30 \sin 3\theta - 1.15 \sin 5\theta + 3.37 \sin 7\theta \\ &= 13.12 \cos (\theta - 0.3^\circ) - 5.04 \cos (3\theta - 3.3^\circ) \\ &- 18.76 \cos (5\theta - 3.6^\circ) + 19.59 \cos (7\theta - 9.9^\circ). \end{aligned} \right\} \quad (14)$$

96. The effective value of this current is given as the square root of the sum of squares of the effective values of the individual harmonics, thus:

$$I = \sqrt{\sum \frac{a^2}{2} + \sum \frac{b^2}{2}} = 21.6 \text{ amp.}$$

As the voltage between line and neutral is 25,400 effective, this gives $Q = 25,400 \times 21.6 = 540,000$ volt-amperes, or 540 kv.-amp. per line, thus a total of $3Q = 1620$ kv.-amp. charging current of the transmission line, when using the e.m.f. wave of these old generators.

It thus would require a minimum of 3 of the 750-kw. generators to keep the voltage on the line, even if no power whatever is delivered from the line.

If the supply voltage of the transmission line were a perfect sine wave, it would, at 44,000 volts between the lines, be given by

$$e_1 = 36 \sin \theta, \quad \dots \quad (15)$$

which is the fundamental, or first harmonic, of equation (9).

Then the current i would also be a sine wave, and would be given by

$$\left. \begin{aligned} i_1 &= a_1 \cos \theta + b_1 \sin \theta, \\ &= 13.12 \cos \theta + 0.07 \sin \theta, \\ &= 13.12 \cos (\theta - 0.3^\circ), \end{aligned} \right\} \quad \dots \quad (16)$$

and its effective value would be

$$I_1 = \frac{13.12}{\sqrt{2}} = 9.3 \text{ amp.} \quad \dots \quad (17)$$

This would correspond to a kv.-amp. input to the line

$$3Q_1 = 3 \times 25.4 \times 9.3 = 710 \text{ kv.-amp.}$$

The distortion of the voltage wave, as given by equation (1), thus increases the charging volt-amperes of the line from 710

kv.-amp. to 1620 kv.-amp. or 2.28 times, and while with a sine wave of voltage, one of the 750-kw. generators would easily be able to supply the charging current of the line, due to the

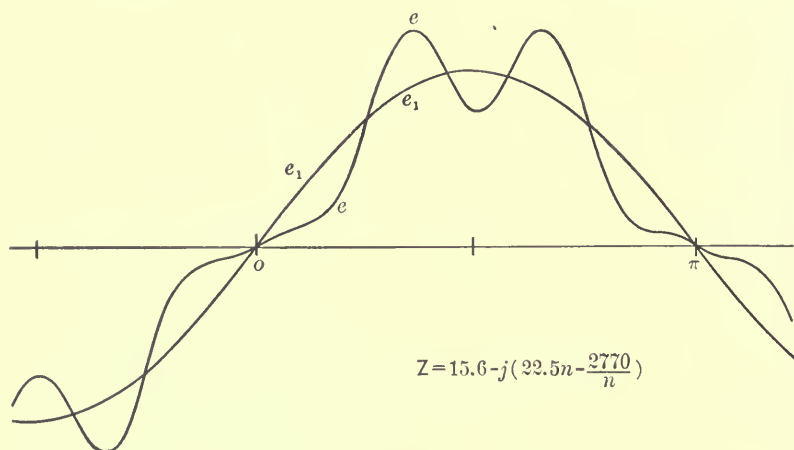


FIG. 47.

wave shape distortion, more than two generators are required. It would, therefore, not be economical to use these generators on the transmission line, if they can be used for any other purposes, as short-distance distribution.

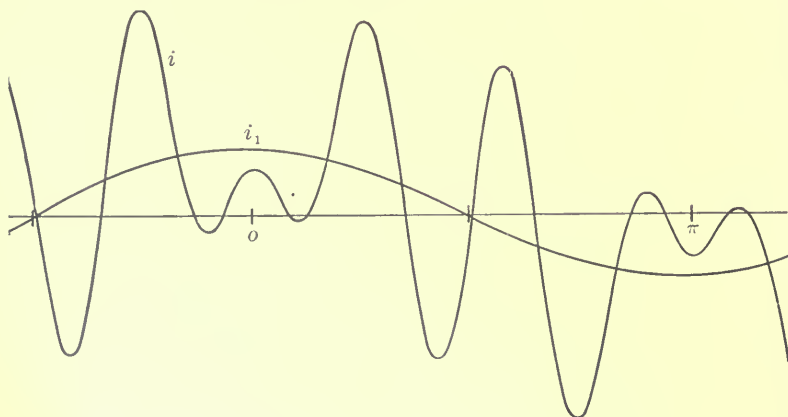


FIG. 48.

In Figs. 47 and 48 are plotted the voltage wave and the current wave, from equations (9) and (14) respectively, and

the numerical values, from 10 deg. to 10 deg., recorded in Table XII.

In Figs. 47 and 48 the fundamental sine wave of voltage and current are also shown. As seen, the distortion of current is enormous, and the higher harmonics predominate over the fundamental. Such waves are occasionally observed as charging currents of transmission lines or cable systems.

97. Assuming now that a reactive coil is inserted in series with the transmission line, between the step-up transformers and the line, what will be the voltage at the terminals of this reactive coil, with the distorted wave of charging current traversing the reactive coil, and how does it compare with the voltage existing with a sine wave of charging current?

Let L =inductance, thus $x=2\pi fL$ =reactance of the coil, and neglecting its resistance, the voltage at the terminals of the reactive coil is given by

$$e' = -x \frac{di}{dt}, \quad . \quad . \quad . \quad . \quad . \quad . \quad (18)$$

Substituting herein the equation of current, (11), gives

$$e' = x \left\{ a_1 \sin \theta + 3a_3 \sin 3\theta + 5a_5 \sin 5\theta + 7a_7 \sin 7\theta \right. \\ \left. - b_1 \cos \theta - 3b_3 \cos 3\theta - 5b_5 \cos 5\theta - 7b_7 \cos 7\theta \right\}; \quad (19)$$

hence, substituting the numerical values (13),

$$e' = x \left\{ 13.12 \sin \theta - 15.09 \sin 3\theta - 93.6 \sin 5\theta + 135.1 \sin 7\theta \right. \\ \left. - 0.07 \cos \theta + 0.90 \cos 3\theta + 5.75 \cos 5\theta - 23.6 \cos 7\theta \right\} \\ = x \left\{ 13.12 \sin (\theta - 0.3^\circ) - 15.12 \sin (3\theta - 3.3^\circ) \right. \\ \left. - 93.8 \sin (5\theta - 3.6^\circ) + 139.1 \sin (7\theta - 9.9^\circ) \right\}; \quad (20)$$

This voltage gives the effective value

$$E' = x \sqrt{\frac{1}{2} \{ 13.12^2 + 15.12^2 + 93.8^2 + 139.1^2 \}} = 119.4x,$$

while the effective value with a sine wave would be from (17),

$$E_1' = xI_1 = 9.3x;$$

hence, the voltage across the reactance x has been increased 12.8 times by the wave distortion.

The instantaneous values of the voltage e' are given in the last column of Table XII, and plotted in Fig. 49, for $x=1$. As seen from Fig. 49, the fundamental wave has practically

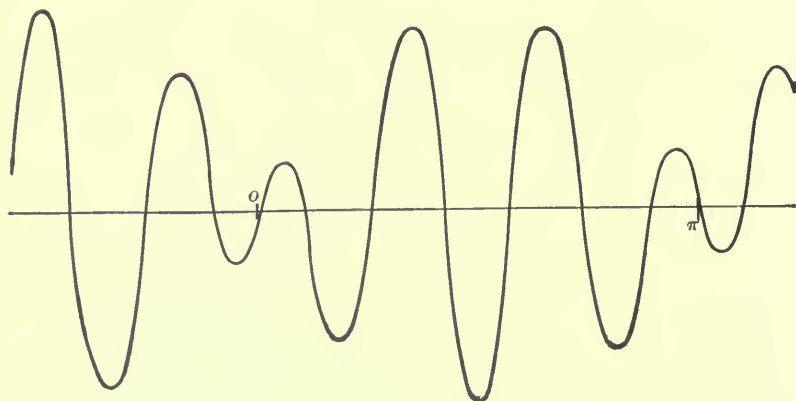


FIG. 49.

vanished, and the voltage wave is the seventh harmonic, modified by the fifth harmonic.

TABLE XII

θ	e	i	e'	θ	e	i	e'
0	-0.10	+ 8.67	- 17	90	27.41	- 4.15	-200
10	+2.23	+ 5.30	+ 46	100	31.77	+26.19	-106
20	3.74	- 0.86	+ 3	110	40.57	+24.99	+119
30	7.47	+ 7.39	+131	120	42.70	- 8.10	+182
40	17.35	+30.39	-116	130	33.14	-38.79	+ 93
50	31.70	+38.58	+ 36	140	18.03	-36.65	- 96
60	42.06	+15.66	+167	150	6.99	-13.41	-138
70	40.33	-19.01	+159	160	2.88	+ 2.43	- 31
80	32.87	-29.13	- 54	170	1.97	- 1.00	+ 54
90	27.41	- 4.15	-200	180	+0.10	- 8.67	+ 17

CHAPTER IV.

MAXIMA AND MINIMA.

98. In engineering investigations the problem of determining the maxima and the minima, that is, the extrema of a function, frequently occurs. For instance, the output of an electric machine is to be found, at which its efficiency is a maximum, or, it is desired to determine that load on an induction motor which gives the highest power-factor; or, that voltage

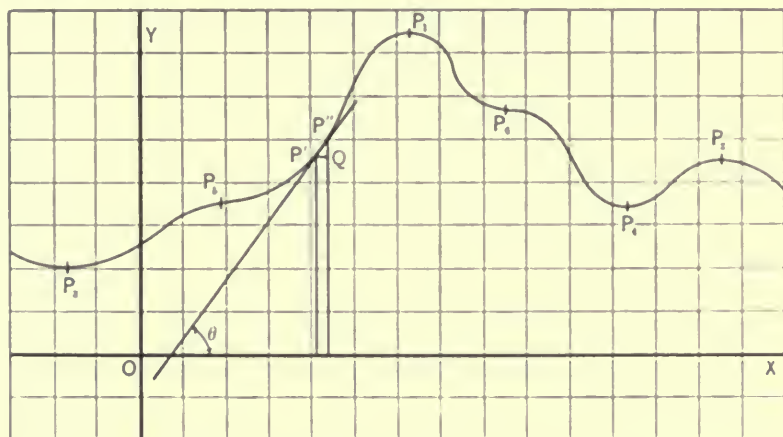


FIG. 50. Graphic Solution of Maxima and Minima.

which makes the cost of a transmission line a minimum; or, that speed of a steam turbine which gives the lowest specific steam consumption, etc.

The maxima and minima of a function, $y = f(x)$, can be found by plotting the function as a curve and taking from the curve the values x , y , which give a maximum or a minimum. For instance, in the curve Fig. 50, maxima are at P_1 and P_2 , minima at P_3 and P_4 . This method of determining the extrema of functions is necessary, if the mathematical expression between

x and y , that is, the function $y=f(x)$, is unknown, or if the function $y=f(x)$ is so complicated, as to make the mathematical calculation of the extrema impracticable. As examples of this method the following may be chosen:

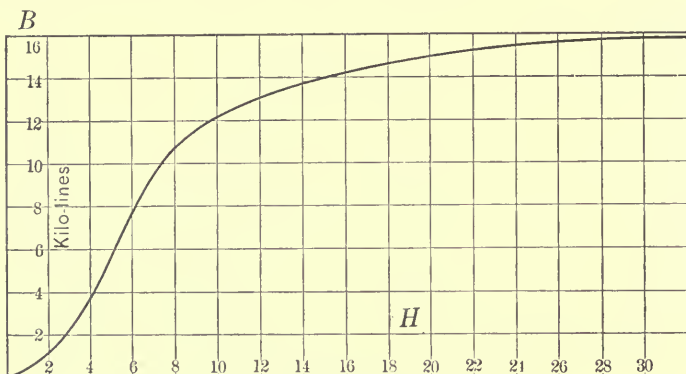


FIG. 51. Magnetization Curve.

Example 1. Determine that magnetic density B , at which the permeability μ of a sample of iron is a maximum. The relation between magnetic field intensity H , magnetic density B and permeability μ cannot be expressed in a mathematical equation, and is therefore usually given in the form of an

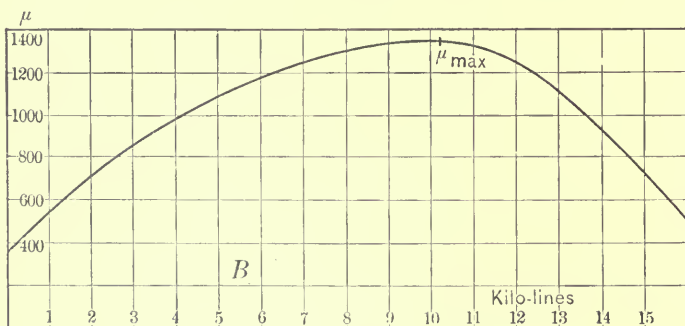


FIG. 52. Permeability Curve.

empirical curve, relating B and H , as shown in Fig. 51. From this curve, corresponding values of B and H are taken, and their ratio, that is, the permeability $\mu = \frac{B}{H}$, plotted against B as abscissa. This is done in Fig. 52. Fig. 52 then shows that a maximum

occurs at point $\mu_{\max}$, for $B=10.2$ kilolines, $\mu=1340$, and minima at the starting-point P_2 , for $B=0$, $\mu=370$, and also for $B=\infty$, where by extrapolation $\mu=1$.

Example 2. Find that output of an induction motor which gives the highest power-factor. While theoretically an equation can be found relating output and power-factor of an induction motor, the equation is too complicated for use. The most convenient way of calculating induction motors is to calculate in tabular form for different values of slip s , the torque, output, current, power and volt-ampere input, efficiency, power-factor, etc., as is explained in "Theoretical Elements of Electrical Engineering," third edition, p. 363. From this

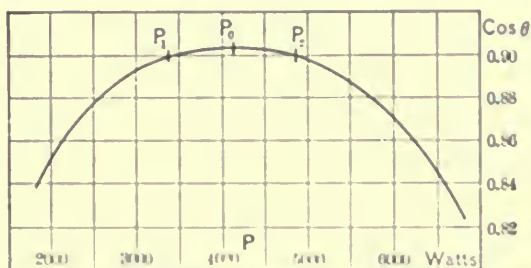


Fig. 53. Power-factor Maximum of Induction Motor.

table corresponding values of power output P and power-factor $\cos \theta$ are taken and plotted in a curve, Fig. 53, and the maximum derived from this curve is $P=4120$, $\cos \theta=0.904$.

For the purpose of determining the maximum, obviously not the entire curve needs to be calculated, but only a short range near the maximum. This is located by trial. Thus in the present instance, P and $\cos \theta$ are calculated for $s=0.1$ and $s=0.2$. As the latter gives lower power-factor, the maximum power-factor is below $s=0.2$. Then $s=0.05$ is calculated and gives a higher value of $\cos \theta$ than $s=0.1$; that is, the maximum is below $s=0.1$. Then $s=0.02$ is calculated, and gives a lower value of $\cos \theta$ than $s=0.05$. The maximum value of $\cos \theta$ thus lies between $s=0.02$ and $s=0.1$, and only the part of the curve between $s=0.02$ and $s=0.1$ needs to be calculated for the determination of the maximum of $\cos \theta$, as is done in Fig. 53.

99. When determining an extremum of a function $y=f(x)$, by plotting it as a curve, the value of x , at which the extreme

occurs, is more or less inaccurate, since at the extreme the curve is horizontal. For instance, in Fig. 53, the maximum of the curve is so flat that the value of power P , for which $\cos \theta$ became a maximum, may be anywhere between $P=4000$ and $P=4300$, within the accuracy of the curve.

In such a case, a higher accuracy can frequently be reached by not attempting to locate the exact extreme, but two points of the same ordinate, on each side of the extreme. Thus in Fig. 53 the power P_0 , at which the maximum power factor $\cos \theta=0.904$ is reached, is somewhat uncertain. The value of power-factor, somewhat below the maximum, $\cos \theta=0.90$, is reached before the maximum, at $P_1=3400$, and after the maximum, at $P_2=4840$. The maximum then may be calculated as half-way between P_1 and P_2 , that is, at $P_0=\frac{1}{2}\{P_1+P_2\}=4120$ watts.

This method gives usually more accurate results, but is based on the assumption that the curve is symmetrical on both sides of the extreme, that is, falls off from the extreme value at the same rate for lower as for higher values of the abscissas. Where this is not the case, this method of interpolation does not give the exact maximum.

Example 3. The efficiency of a steam turbine nozzle, that is, the ratio of the kinetic energy of the steam jet to the energy of the steam available between the two pressures between which the nozzle operates, is given in Fig. 54, as determined by experiment. As abscissas are used the nozzle mouth opening, that is, the widest part of the nozzle at the exhaust end, as fraction of that corresponding to the exhaust pressure, while the nozzle throat, that is, the narrowest part of the nozzle, is assumed as constant. As ordinates are plotted the efficiencies. This curve is not symmetrical, but falls off from the maximum, on the sides of larger nozzle mouth, far more rapidly than on the side of smaller nozzle mouth. The reason is that with too large a nozzle mouth the expansion in the nozzle is carried below the exhaust pressure p_2 , and steam eddies are produced by this overexpansion.

The maximum efficiency of 94.6 per cent is found at the point P_0 , at which the nozzle mouth corresponds to the exhaust pressure. If, however, the maximum is determined as mid-way between two points P_1 and P_2 , on each side of the maximum,

at which the efficiency is the same, 93 per cent, a point P_0' is obtained, which lies on one side of the maximum.

With unsymmetrical curves, the method of interpolation thus does not give the exact extreme. For most engineering purposes this is rather an advantage. The purpose of determining the extreme usually is to select the most favorable operating conditions. Since, however, in practice the operating conditions never remain perfectly constant, but vary to some extent, the most favorable operating condition in Fig. 54 is not that where the average value gives the maximum efficiency

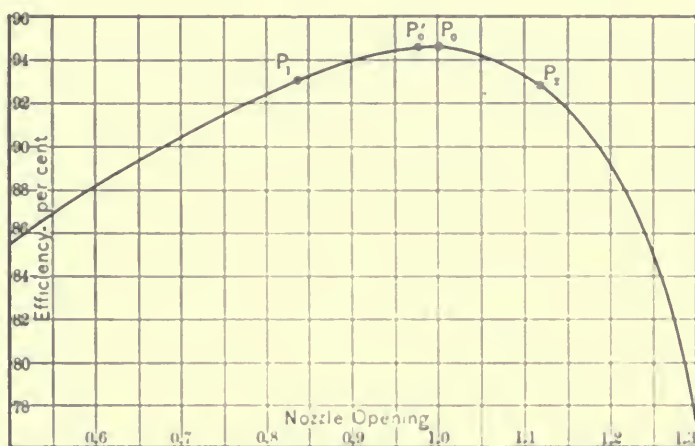


FIG. 54. Steam Turbine Nozzle Efficiency; Determination of Maximum.

(point P_0), but the most favorable operating condition is that, where the average efficiency during the range of pressure, occurring in operation, is a maximum.

If the steam pressure, and thereby the required expansion ratio, that is, the theoretically correct size of nozzle mouth, should vary during operation by 25 per cent from the average, when choosing the maximum efficiency point P_0 as average, the efficiency during operation varies on the part of the curve between P_1 (91.4 per cent) and P_2 (85.2 per cent), thus averaging lower than by choosing the point P_0' (6.25 per cent below P_0) as average. In the latter case, the efficiency varies on the part of the curve from the P_1' (90.1 per cent) to P_2' (90.1 per cent). (Fig. 55.)

Thus in apparatus design, when determining extrema of a function $y=f(x)$, to select them as operating condition, consideration must be given to the shape of the curve, and where the curve is unsymmetrical, the most efficient operating point may not lie at the extreme, but on that side of it at which the curve falls off slower, the more so the greater the range of variation is, which may occur during operation. This is not always realized.

100. If the function $y=f(x)$ is plotted as a curve, Fig. 50, at the extremes of the function, the points P_1, P_2, P_3, P_4 of curve Fig. 50, the tangent on the curve is horizontal, since

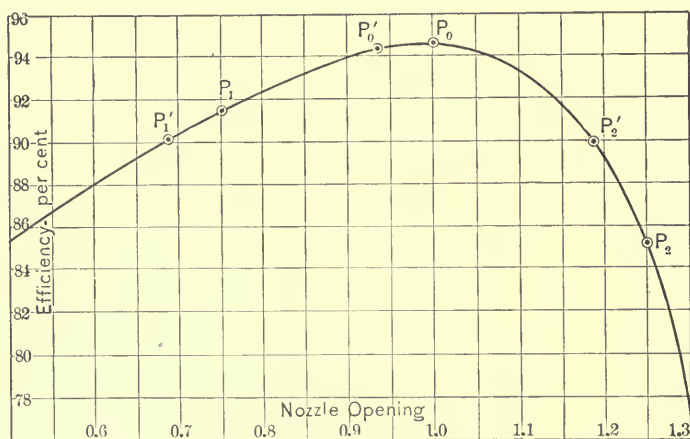


FIG. 55. Steam Turbine Nozzle Efficiency; Determination of Maximum.

at the extreme the function changes from rising to decreasing (maximum, P_1 and P_2), or from decreasing to increasing (minimum, P_3 and P_4), and therefore for a moment passes through the horizontal direction.

In general, the tangent of a curve, as that in Fig. 50, is the line which connects two points P' and P'' of the curve, which are infinitely close together, and, as seen in Fig. 50, the angle θ , which this tangent $P'P''$ makes with the horizontal or X -axis, thus is given by:

$$\tan \theta = \frac{P''Q}{P'Q} = \frac{dy}{dx}.$$

At the extreme, the tangent on the curve is horizontal, that is, $\angle \theta = 0$, and, therefore, it follows that at an extreme of the function,

$$y = f(x), \quad \dots \dots \dots (1)$$

$$\frac{dy}{dx} = 0 \quad \dots \dots \dots (2)$$

The reverse, however, is not necessarily the case; that is, if at a point $x, y : \frac{dy}{dx} = 0$, this point may not be an extreme; that is, a maximum or minimum, but may be a horizontal inflection point, as points P_3 and P_6 are in Fig. 50.

With increasing x , when passing a maximum (P_1 and P_2 , Fig. 50), y rises, then stops rising, and then decreases again. When passing a minimum (P_3 and P_4) y decreases, then stops decreasing, and then increases again. When passing a horizontal inflection point, y rises, then stops rising, and then starts rising again, at P_5 , or y decreases, then stops decreasing, but then starts decreasing again (at P_6).

The points of the function $y = f(x)$, determined by the condition, $\frac{dy}{dx} = 0$, thus require further investigation, whether they represent a maximum, or a minimum, or merely a horizontal inflection point.

This can be done mathematically: for increasing x , when passing a maximum, $\tan \theta$ changes from positive to negative; that is, decreases, or in other words, $\frac{d}{dx} (\tan \theta) < 0$. Since $\tan \theta = \frac{dy}{dx}$, it thus follows that at a maximum $\frac{d^2y}{dx^2} < 0$. Inversely, at a minimum $\tan \theta$ changes from negative to positive, hence increases, that is, $\frac{d}{dx} (\tan \theta) > 0$; or, $\frac{d^2y}{dx^2} > 0$. When passing a horizontal inflection point $\tan \theta$ first decreases to zero at the inflection point, and then increases again; or, inversely, $\tan \theta$ first increases, and then decreases again, that is, $\tan \theta = \frac{dy}{dx}$ has a maximum or a minimum at the inflection point, and therefore, $\frac{d}{dx} (\tan \theta) = \frac{d^2y}{dx^2} = 0$ at the inflection point.

In engineering problems the investigation, whether the solution of the condition of extremes, $\frac{dy}{dx}=0$, represents a minimum, or a maximum, or an inflection point, is rarely required, but it is almost always obvious from the nature of the problem whether a maximum or a minimum occurs, or neither.

For instance, if the problem is to determine the speed at which the efficiency of a motor is a maximum, the solution: speed=0, obviously is not a maximum but a minimum, as at zero speed the efficiency is zero. If the problem is, to find the current at which the output of an alternator is a maximum, the solution $i=0$ obviously is a minimum, and of the other two solutions, i_1 and i_2 , the larger value, i_2 , again gives a minimum, zero output at short-circuit current, while the intermediate value i_1 gives the maximum.

101. The extremes of a function, therefore, are determined by equating its differential quotient to zero, as is illustrated by the following examples:

Example 4. In an impulse turbine, the speed of the jet (steam jet or water jet) is S_1 . At what peripheral speed S_2 is the output a maximum.

The impulse force is proportional to the relative speed of the jet and the rotating impulse wheel; that is, to (S_1-S_2) . The power is impulse force times speed S_2 ; hence,

$$P = kS_2(S_1 - S_2), \quad . \quad . \quad . \quad . \quad . \quad . \quad (3)$$

and is an extreme for the value of S_2 , given by $\frac{dP}{dS_2}=0$; hence,

$$S_1 - 2S_2 = 0 \quad \text{and} \quad S_2 = \frac{S_1}{2}; \quad . \quad . \quad . \quad . \quad . \quad (4)$$

that is, when the peripheral speed of the impulse wheel equals half the jet velocity.

Example 5. In a transformer of constant impressed e.m.f. $e_0=2300$ volts; the constant loss, that is, loss which is independent of the output (iron loss), is $P_i=500$ watts. The internal resistance (primary and secondary combined) is $r=20$

ohms. At what current i is the efficiency of the transformer a maximum; that is, the percentage loss, λ , a minimum?

$$\text{The loss is } P = P_i + ri^2 = 500 + 20i^2. \quad (5)$$

$$\text{The power input is } P_1 = ei = 2300i; \quad (6)$$

hence, the percentage loss is,

$$\lambda = \frac{P}{P_1} = \frac{P_i + ri^2}{ei}, \quad (7)$$

and this is an extreme for the value of current i , given by

$$\frac{d\lambda}{di} = 0;$$

hence,

$$\frac{(P_i + ri^2)e - ei(2ri)}{e^2i^2} = 0;$$

or,

$$P_i - ri^2 = 0 \quad \text{and} \quad i = \sqrt{\frac{P_i}{r}} = 5 \text{ amperes}, \quad (8)$$

and the output is $P_o = ei = 11,500$ watts. The loss is, $P - P_i + ri^2 = 2P_i = 1000$ watts; that is, the i^2r loss or variable loss, is equal to the constant loss P_i . The percentage loss is,

$$\lambda = \frac{P}{P_i} = \frac{\sqrt{P_i r}}{e} = 0.087 = 8.7 \text{ per cent},$$

and the maximum efficiency thus is,

$$1 - \lambda = 0.913 = 91.3 \text{ per cent}.$$

102. Usually, when the problem is given, to determine those values of x for which y is an extreme, y cannot be expressed directly as function of x , $y = f(x)$, as was done in examples (4) and (5), but y is expressed as function of some other quantities, $y = f(u, v, \dots)$, and then equations between $u, v, \dots$ and x are found from the conditions of the problem, by which expressions of x are substituted for $u, v, \dots$, as shown in the following example:

Example 6. There is a constant current i_0 through a circuit containing a resistor of resistance r_0 . This resistor r_0

is shunted by a resistor of resistance r . What must be the resistance of this shunting resistor r , to make the power consumed in r , a maximum? (Fig. 56.)

Let i be the current in the shunting resistor r . The power consumed in r then is,

$$P = ri^2 \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (9)$$

The current in the resistor r_0 is $i_0 - i$, and therefore the voltage consumed by r_0 is $r_0(i_0 - i)$, and the voltage consumed by r is ri , and as these two voltages must be equal, since both

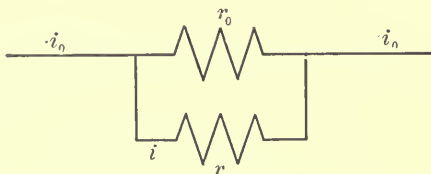


FIG. 56. Shunted Resistor.

resistors are in shunt with each other, thus receive the same voltage,

$$ri = r_0(i_0 - i),$$

and, herefrom, it follows that,

$$i = \frac{r_0}{r + r_0} i_0. \quad . \quad . \quad . \quad . \quad . \quad . \quad (10)$$

Substituting this in equation (9) gives,

$$P = \frac{rr_0^2 i_0^2}{(r + r_0)^2}, \quad . \quad . \quad . \quad . \quad . \quad . \quad (11)$$

and this power is an extreme for $\frac{dP}{dr} = 0$; hence:

$$\frac{(r + r_0)^2 r_0^2 i_0^2 - rr_0^2 i_0^2 2(r + r_0)}{(r + r_0)^4} = 0;$$

hence,

$$r = r_0; \quad . \quad . \quad . \quad . \quad . \quad . \quad (12)$$

that is, the power consumed in r is a maximum, if the resistor r of the shunt equals the resistance r_0 .

The current in r then is, by equation (10),

$$i = \frac{i_0}{2},$$

and the power is,

$$P = \frac{r_0 i_0^2}{4}.$$

103. If, after the function $y=f(x)$ (the equation (11) in example (6)) has been derived, the differentiation $\frac{dy}{dx}=0$ is immediately carried out, the calculation is very frequently much more complicated than necessary. It is, therefore, advisable *not* to differentiate immediately, but first to simplify the function $y=f(x)$.

If y is an extreme, any expression differing therefrom by a constant term, or a constant factor, etc., also is an extreme. So also is the reciprocal of y , or its square, or square root, etc.

Thus, before differentiation, constant terms and constant factors can be dropped, fractions inverted, the expression raised to any power or any root thereof taken, etc.

For instance, in the preceding example, in equation (11),

$$P = \frac{r r_0^2 i_0^2}{(r + r_0)^2},$$

the value of r is to be found, which makes P a maximum.

If P is an extreme,

$$y_1 = \frac{r}{(r + r_0)^2},$$

which differs from P by the omission of the constant factor $r_0^2 i_0^2$, also is an extreme.

The reverse of y_1 ,

$$y_2 = \frac{(r + r_0)^2}{r},$$

is also an extreme. (y_2 is a minimum, where y_1 is a maximum, and inversely.)

Therefore, the equation (11) can be simplified to the form:

$$y_2 = \frac{(r + r_0)^2}{r} = r + 2r_0 + \frac{r_0^2}{r},$$

and, leaving out the constant term $2r_0$, gives the final form,

$$y_3 = r + \frac{r_0^2}{r}. \quad \dots \quad (13)$$

This differentiated gives,

$$\frac{dy_3}{dr} = 1 - \frac{r_0^2}{r^2} = 0;$$

hence,

$$r = r_0.$$

104. Example 7. From a source of constant alternating e.m.f. e , power is transmitted over a line of resistance r_0 and reactance x_0 into a non-inductive load. What must be the resistance r of this load to give maximum power?

If i = current transmitted over the line, the power delivered at the load of resistance r is

$$P = ri^2. \quad \dots \quad (14)$$

The total resistance of the circuit is $r + r_0$; the reactance is x_0 ; hence the current is

$$i = \frac{e}{\sqrt{(r + r_0)^2 + x_0^2}}, \quad \dots \quad (15)$$

and, by substituting in equation (14), the power is

$$P = \frac{re^2}{(r + r_0)^2 + x_0^2}, \quad \dots \quad (16)$$

if P is an extreme, by omitting e^2 and inverting,

$$\begin{aligned} y_1 &= \frac{(r + r_0)^2 + x_0^2}{r} \\ &= r + 2r_0 + \frac{r_0^2 + x_0^2}{r}, \end{aligned}$$

is also an extreme, and likewise,

$$y_2 = r + \frac{r_0^2 + x_0^2}{r},$$

is an extreme.

Differentiating, gives:

$$\frac{dy_2}{dr} = 1 - \frac{r_0^2 + x_0^2}{r^2} = 0,$$

$$r = \sqrt{r_0^2 + x_0^2}. \quad . \quad . \quad . \quad . \quad . \quad . \quad (17)$$

Wherefrom follows, by substituting in equation (16),

$$P = \frac{\sqrt{r_0^2 + x_0^2} e^2}{(r_0 + \sqrt{r_0^2 + x_0^2})^2 + x_0^2}$$

$$= \frac{e^2}{2(r_0 + \sqrt{r_0^2 + x_0^2})}. \quad . \quad . \quad . \quad . \quad . \quad (18)$$

Very often the function $y=f(x)$ can by such algebraic operations, which do not change an extreme, be simplified to such an extent that differentiation becomes entirely unnecessary, but the extreme is immediately seen; the following example will serve to illustrate:

Example 8. In the same transmission circuit as in example (7), for what value of r is the current i a maximum?

The current i is given, by equation (15),

$$i = \frac{e}{\sqrt{(r + r_0)^2 + x_0^2}}.$$

Dropping e and reversing, gives,

$$y_1 = \sqrt{(r + r_0)^2 + x_0^2};$$

Squaring, gives,

$$y_2 = (r + r_0)^2 + x_0^2;$$

dropping the constant term x_0^2 gives

$$y_3 = (r + r_0)^2; \quad . \quad . \quad . \quad . \quad . \quad . \quad (19)$$

taking the square root gives

$$y_4 = r + r_0;$$

immediately differentiated, gives as condition of the extremes:

$$\frac{di}{dr} = -\frac{2(r+r_0)}{2\{(r+r_0)^2+x_0^2\}\sqrt{(r+r_0)^2+x_0^2}} = 0;$$

hence, either $r+r_0=0;$ (22)

or, $(r+r_0)^2+x_0^2=\infty$ (23)

the latter equation gives $r=\infty$; hence $i=0$, the minimum value of current.

The former equation gives

$$r=-r_0, \text{ (24)}$$

as the value of the resistance, which gives maximum current, and the current would then be, by substituting (24) into (15),

$$i = \frac{e}{x_0}. \text{ (25)}$$

The solution (24), however, has no engineering meaning, as the resistance r cannot be negative.

Hence, mathematically, there exists no maximum value of i in the range of r which can occur in engineering, that is, within the range, $0 < r < \infty$.

In such a case, where the extreme falls outside of the range of numerical values, to which the engineering quantity is limited, it follows that within the engineering range the quantity continuously increases toward one limit and continuously decreases toward the other limit, and that therefore the two limits of the engineering range of the quantity give extremes. Thus $r=0$ gives the maximum, $r=\infty$ the minimum of current.

106. Example 10. An alternating-current generator, of generated e.m.f. $e=2500$ volts, internal resistance $r_0=0.25$ ohms, and synchronous reactance $x_0=10$ ohms, is loaded by a circuit comprising a resistor of constant resistance $r=20$ ohms, and a reactor of reactance x in series with the resistor r . What value of reactance x gives maximum output?

If i =current of the alternator, its power output is

$$P=ri^2=20i^2; \text{ (26)}$$

the total resistance is $r + r_0 = 20.25$ ohms; the total reactance is $x + x_0 = 10 + x$ ohms, and therefore the current is

$$i = \frac{e}{\sqrt{(r + r_0)^2 + (x + x_0)^2}}, \quad \dots \quad (27)$$

and the power output, by substituting (27) in (26), is

$$P = \frac{re^2}{(r + r_0)^2 + (x + x_0)^2} = \frac{20 \times 2500^2}{20.25^2 + (10 + x)^2}, \quad \dots \quad (28)$$

Simplified, this gives

$$y_1 = (r + r_0)^2 + (x + x_0)^2; \quad \dots \quad (29)$$

$$y_2 = (x + x_0)^2;$$

hence,

$$\frac{dy_2}{dx} = 2(x + x_0) = 0;$$

and

$$x = -x_0 = -10 \text{ ohms}; \quad \dots \quad (30)$$

that is, a negative, or condensive reactance of 10 ohms. The power output would then be, by substituting (30) into (28),

$$P = \frac{re^2}{(r + r_0)^2} = \frac{20 + 2500^2}{20.25^2} \text{ watts} = 305 \text{ kw.} \quad \dots \quad (31)$$

If, however, a condensive reactance is excluded, that is, it is assumed that $x > 0$, no mathematical extreme exists in the range of the variable x , which is permissible, and the extreme is at the end of the range, $x = 0$, and gives

$$P = \frac{re_0^2}{(r + r_0)^2 + x_0^2} = 245 \text{ kw.} \quad \dots \quad (32)$$

107. Example 11. In a 500-kw. alternator, at voltage $e = 2500$, the friction and windage loss is $P_w = 6$ kw., the iron loss $P_i = 24$ kw., the field excitation loss is $P_f = 6$ kw., and the armature resistance $r = 0.1$ ohm. At what load is the efficiency a maximum?

The sum of the losses is:

$$P = P_w + P_i + P_f + ri^2 = 36,000 + 0.1i^2. \quad (33)$$

The output is

$$P_0 = ei = 2500i; \quad (34)$$

hence, the efficiency is

$$\eta = \frac{P_0}{P_0 + P} = \frac{ei}{ei + P_w + P_i + P_f + ri^2} = \frac{2500i}{36000 + 2500i + 0.1i^2}; \quad (35)$$

or, simplified,

$$y_1 = \frac{P_w + P_i + P_f}{i} + ri;$$

hence,

$$\frac{dy_1}{di} = r - \frac{P_w + P_i + P_f}{i^2},$$

and,

$$i = \sqrt{\frac{P_w + P_i + P_f}{r}} = \sqrt{\frac{36000}{0.1}} = 600 \text{ amperes.} \quad (36)$$

and the output, at which the maximum efficiency occurs, by substituting (36) into (34), is

$$P = ei = 1500 \text{ kw.},$$

that is, at three times full load.

Therefore, this value is of no engineering importance, but means that at full load and at all practical overloads the maximum efficiency is not yet reached, but the efficiency is still rising.

108. Frequently in engineering calculations extremes of engineering quantities are to be determined, which are functions of two or more independent variables. For instance, the maximum power is required which can be delivered over a transmission line into a circuit, in which the resistance as well as the reactance can be varied independently. In other words, if

$$y = f(u, v) \quad (37)$$

is a function of two independent variables u and v , such a pair of values of u and of v is to be found, which makes y a maximum, or minimum.

Choosing any value u_0 , of the independent variable u , then a value of v can be found, which gives the maximum (or minimum) value of y , which can be reached for $u=u_0$. This is done by differentiating $y=f(u_0, v)$, over v , thus:

$$\frac{df(u_0, v)}{dv} = 0, \quad . \quad . \quad . \quad . \quad . \quad . \quad (38)$$

From this equation (38), a value,

$$v=f_1(u_0), \quad . \quad . \quad . \quad . \quad . \quad . \quad (39)$$

is derived, which gives the maximum value of y , for the given value of u_0 , and by substituting (39) into (38),

$$y=f_2(u_0), \quad . \quad . \quad . \quad . \quad . \quad . \quad (40)$$

is obtained as the equation, which relates the different extremes of y , that correspond to the different values of u_0 , with u_0 .

Herefrom, then, that value of u_0 is found which gives the maximum of the maxima, by differentiation:

$$\frac{df_2(u_0)}{du_0} = 0. \quad . \quad . \quad . \quad . \quad . \quad . \quad (41)$$

Geometrically, $y=f(u, v)$ may be represented by a surface in space, with the coordinates y, u, v . $y=f(u_0, v)$, then, represents the curve of intersection of this surface with the plane $u_0 = \text{constant}$, and the differentiation gives the maximum point of this intersection curve. $y=f_2(u_0)$ then gives the curve in space, which connects all the maxima of the various intersections with the u_0 planes, and the second differentiation gives the maximum of this maximum curve $y=f_2(u_0)$, or the maximum of the maxima (or more correctly, the extreme of the extremes).

Inversely, it is possible first to differentiate over u , thus,

$$\frac{df(u, v_0)}{du} = 0, \quad . \quad . \quad . \quad . \quad . \quad . \quad (42)$$

and thereby get

$$u=f_3(v_0). \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (43)$$

as the value of u , which makes y a maximum for the given value of $v=v_0$, and substituting (43) into (42),

$$y=f_4(v_0), \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (44)$$

is obtained as the equation of the maxima, which differentiated over v_0 , thus,

$$\frac{df_4(v_0)}{dv_0}=0. \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (45)$$

gives the maximum of the maxima.

Geometrically, this represents the consideration of the intersection curves of the surface with the planes $v=\text{constant}$.

However, equations (38) and (41) (respectively (42) and (45)) give an extremum only, if both equations represent maxima, or both minima. If one of the equations represents a maximum, the other a minimum, the point is not an extremum, but a saddle point, so called from the shape of the surface $y=f(u, v)$ near this point.

The working of this will be plain from the following example:

109. Example 12. The alternating voltage $e=30,000$ is impressed upon a transmission line of resistance $r_0=20$ ohms and reactance $x_0=50$ ohms.

What should be the resistance r and the reactance x of the receiving circuit to deliver maximum power?

Let i =current delivered into the receiving circuit. The total resistance is $(r+r_0)$; the total reactance is $(x+x_0)$; hence, the current is

$$i=\frac{e}{\sqrt{(r+r_0)^2+(x+x_0)^2}}. \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (46)$$

$$\text{The power output is} \quad P=ri^2; \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (47)$$

hence, substituting (46), gives

$$P=\frac{re^2}{(r+r_0)^2+(x+x_0)^2}. \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (48)$$

(a) For any given value of r , the reactance x , which gives maximum power, is derived by $\frac{dP}{dx}=0$.

P simplified, gives $y_1 = (x + x_0)^2$; hence,

$$\frac{dy_1}{dx} = 2(x + x_0) = 0 \quad \text{and} \quad x = -x_0 \quad . \quad . \quad . \quad (49)$$

that is, for any chosen resistance r , the power is a maximum, if the reactance of the receiving circuit is chosen equal to that of the line, but of opposite sign, that is, as condensive reactance.

Substituting (49) into (48) gives the maximum power available for a chosen value of r , as:

$$P_0 = \frac{re^2}{(r + r_0)^2}; \quad . \quad . \quad . \quad . \quad . \quad (50)$$

or, simplified,

$$y_2 = \frac{(r + r_0)^2}{r} \quad \text{and} \quad y_3 = r + \frac{r_0^2}{r};$$

hence,

$$\frac{dy_3}{dr} = 1 - \frac{r_0^2}{r^2} \quad \text{and} \quad r = r_0, \quad . \quad . \quad . \quad (51)$$

and by substituting (51) into (50), the maximum power is,

$$P_{\max} = \frac{e^2}{4r_0}. \quad . \quad . \quad . \quad . \quad . \quad (52)$$

(b) For any given value of x , the resistance r , which gives maximum power, is given by $\frac{dP}{dr} = 0$.

P simplified gives,

$$y_1 = \frac{(r + r_0)^2 + (x + x_0)^2}{r}; \quad y_2 = r + \frac{r_0^2 + (x + x_0)^2}{r};$$

$$\frac{dy_2}{dr} = 1 + \frac{r_0^2 + (x + x_0)^2}{r^2} = 0.$$

$$r = \sqrt{r_0^2 + (x + x_0)^2}, \quad . \quad . \quad . \quad . \quad (53)$$

which is the value of r , that for any given value of x , gives maximum power, and this maximum power by substituting (53) into (48) is,

$$\begin{aligned} P_0 &= \frac{\sqrt{r_0^2 + (x + x_0)^2} e^2}{[r_0 + \sqrt{r_0^2 + (x + x_0)^2}]^2 + (x + x_0)^2} \\ &= \frac{e^2}{2\{r_0 + \sqrt{r_0^2 + (x + x_0)^2}\}}; \quad . \quad . \quad . \quad . \quad (54) \end{aligned}$$

which is the maximum power that can be transmitted into a receiving circuit of reactance x .

The value of x , which makes this maximum power P_0 the highest maximum, is given by $\frac{dP_0}{dx} = 0$.

P_0 simplified gives

$$y_3 = r_0 + \sqrt{r_0^2 + (x + x_0)^2};$$

$$y_4 = \sqrt{r_0^2 + (x + x_0)^2};$$

$$y_5 = r_0^2 + (x + x_0)^2;$$

$$y_6 = (x + x_0)^2;$$

$$y_7 = (x + x_0);$$

and this value is a maximum for $(x + x_0) = 0$; that is, for

$$x = -x_0. \quad (55)$$

NOTE. If x cannot be negative, that is, if only inductive reactance is considered, $x = 0$ gives the maximum power, and the latter then is

$$P_{\max} = \frac{e^2}{2\{r_0 + \sqrt{r_0^2 + x_0^2}\}}, \quad (56)$$

the same value as found in problem (7), equation (18).

Substituting (55) and (54) gives again equation (52), thus,

$$P_{\max} = \frac{e^2}{4r_0}.$$

110. Here again, it requires consideration, whether the solution is practicable within the limitation of engineering constants.

With the numerical constants chosen, it would be

$$P_{\max} = \frac{e^2}{4r_0} = \frac{30000^2}{80} = 11,250 \text{ kw.};$$

$$i = \frac{e}{2r_0} = 750 \text{ amperes.}$$

and the voltage at the receiving end of the line would be

$$e' = i\sqrt{r^2 + x^2} = 750\sqrt{20^2 + 50^2} = 40,400 \text{ volts};$$

that is, the voltage at the receiving end would be far higher than at the generator end, the current excessive, and the efficiency of transmission only 50 per cent. This extreme case thus is hardly practicable, and the conclusion would be that by the use of negative reactance in the receiving circuit, an amount of power could be delivered, at a sacrifice of efficiency, far greater than economical transmission would permit.

In the case, where capacity was excluded from the receiving circuit, the maximum power was given by equation (56) as

$$P_{\max} = \frac{e_0^2}{2\{r_0 + \sqrt{r_0^2 + x_0^2}\}} = 6100 \text{ kw.}$$

III. Extremes of engineering quantities x , y , are usually determined by differentiating the function,

$$y = f(x), \quad . \quad . \quad . \quad . \quad . \quad . \quad (57)$$

and from the equation,

$$\frac{dy}{dx} = 0, \quad . \quad . \quad . \quad . \quad . \quad . \quad (58)$$

deriving the values of x , which make y an extreme.

Occasionally, however, the equation (58) cannot be solved for x , but is either of higher order in x , or a transcendental equation. In this case, equation (58) may be solved by approximation, or preferably, the function,

$$z = \frac{dy}{dx}, \quad . \quad . \quad . \quad . \quad . \quad . \quad (59)$$

is plotted as a curve, the values of x taken, at which $z=0$, that is, at which the curve intersects the X -axis. For instance:

Example 13. The e.m.f. wave of a three-phase alternator, as determined by oscillograph, is represented by the equation,

$$e = 36000\{\sin \theta - 0.12 \sin (3\theta - 2.3^\circ) - 23 \sin (5\theta - 1.5^\circ) + 0.13 \sin (7\theta - 6.2^\circ)\}. \quad . \quad . \quad . \quad . \quad (60)$$

This alternator, connected to a long-distance transmission line, gives the charging current to the line of

$$i = 13.12 \cos (\theta - 0.3^\circ) - 5.04 \cos (3\theta - 3.3^\circ) - 18.76 \cos (5\theta - 3.6^\circ) + 19.59 \cos (7\theta - 9.9^\circ) \quad (61)$$

(see Chapter III, paragraph 95).

What are the extreme values of this current, and at what phase angles θ do they occur?

The phase angle θ , at which the current i reaches an extreme value, is given by the equation

$$\frac{di}{d\theta} = 0. \quad (62)$$

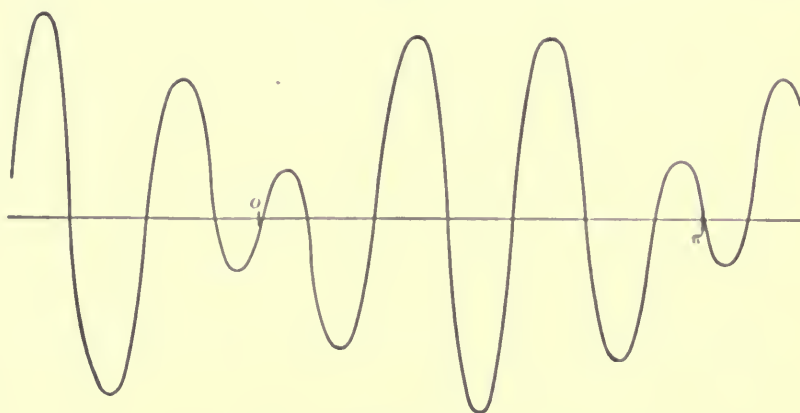


FIG. 57

Substituting (61) into (62) gives,

$$z = \frac{di}{d\theta} = -13.12 \sin (\theta - 0.3^\circ) + 15.12 \sin (3\theta - 3.3^\circ) + 93.8 \sin (5\theta - 3.6^\circ) - 137.1 \sin (7\theta - 9.9^\circ) = 0. \quad (63)$$

This equation cannot be solved for θ . Therefore z is plotted as function of θ by the curve, Fig. 57, and from this curve the values of θ taken at which the curve intersects the zero line. They are:

$$\theta = 1^\circ; 20^\circ; 47^\circ 78'; 101^\circ; 135^\circ; 162^\circ.$$

For these angles θ , the corresponding values of i are calculated by equation (61), and are:

$$i_0 = +9; -1; +39; -30; +30; -42; +4 \text{ amperes.}$$

The current thus has during each period 14 extrema, of which the highest is 42 amperes.

112. In those cases, where the mathematical expression of the function $y=f(x)$ is not known, and the extreme values therefore have to be determined graphically, frequently a greater accuracy can be reached by plotting as a curve the differential of $y=f(x)$ and picking out the zero values instead of plotting $y=f(x)$, and picking out the highest and the lowest points. If the mathematical expression of $y=f(x)$ is not known, obviously the equation of the differential curve $z=\frac{dy}{dx}$ (64) is usually not known either. Approximately, however, it can frequently be plotted from the numerical values of $y=f(x)$, as follows:

If $x_1, x_2, x_3 \dots$ are successive numerical values of x ,
and $y_1, y_2, y_3 \dots$ the corresponding numerical values of y ,
approximate points of the differential curve $z=\frac{dy}{dx}$ are given
by the corresponding values:

$$\text{as abscissas: } \frac{x_2+x_1}{2}; \quad \frac{x_3+x_2}{2}; \quad \frac{x_4+x_3}{2} \dots;$$

$$\text{as ordinates: } \frac{y_2-y_1}{x_2-x_1}; \quad \frac{y_3-y_2}{x_3-x_2}; \quad \frac{y_4-y_3}{x_4-x_3} \dots$$

113. Example 14. In the problem (1), the maximum permeability point of a sample of iron, of which the B, H curve is given as Fig. 51, was determined by taking from Fig. 51 corresponding values of B and H , and plotting $\mu=\frac{B}{H}$, against B in Fig. 52. A considerable inaccuracy exists in this method, in locating the value of B , at which μ is a maximum, due to the flatness of the curve, Fig. 52.

The successive pairs of corresponding values of B and H , as taken from Fig. 51 are given in columns 1 and 2 of Table I.

TABLE I.

B Kilolines,	H	$\mu = \frac{B}{H}$	$J\mu$	B
0	0	370		
1	1.76	570	+200	0.5
2	2.74	730	160	1.5
3	3.47	865	135	2.5
4	4.06	985	120	3.5
5	4.59	1090	105	4.5
6	5.10	1175	85	5.5
7	5.63	1245	70	6.5
8	6.17	1295	50	7.5
9	6.77	1330	35	8.5
10	7.47	1340	+10	9.5
11	8.33	1320	-20	10.5
12	9.60	1250	70	11.5
13	11.60	1120	130	12.5
14	15.10	930	190	13.5
15	20.7	725	205	14.5

In the third column of Table I is given the permeability, $\mu = \frac{B}{H}$, and in the fourth column the increase of permeability, per $B=1$, $J\mu$; the last column then gives the value of B , to which $J\mu$ corresponds.

In Fig. 58, values $J\mu$ are plotted as ordinates, with B as abscissas. This curve passes through zero at $B=9.95$.

The maximum permeability thus occurs at the approximate magnetic density $B=9.95$ kilolines per sq.cm., and not at $B=10.2$, as was given by the less accurate graphical determination of Fig. 52, and the maximum permeability is $\mu_0=1340$.

As seen, the sharpness of the intersection of the differential curve with the zero line permits a far greater accuracy than feasible by the method used in instance (1).

114. As illustration of the method of determining extremes, some further examples are given below:

Example 15. A storage battery of $n=80$ cells is to be connected so as to give maximum power in a constant resistance $r=0.1$ ohm. Each battery cell has the e.m.f. $e_0=2.1$ volts and the internal resistance $r_0=0.02$ ohm. How must the cells be connected?

Assuming the cells are connected with x in parallel, hence $\frac{n}{x}$ in series. The internal resistance of the battery then is

$\frac{n}{x} r_0 = \frac{nr_0}{x^2}$ ohms, and the total resistance of the circuit is $\frac{nr_0}{x^2} + r$.

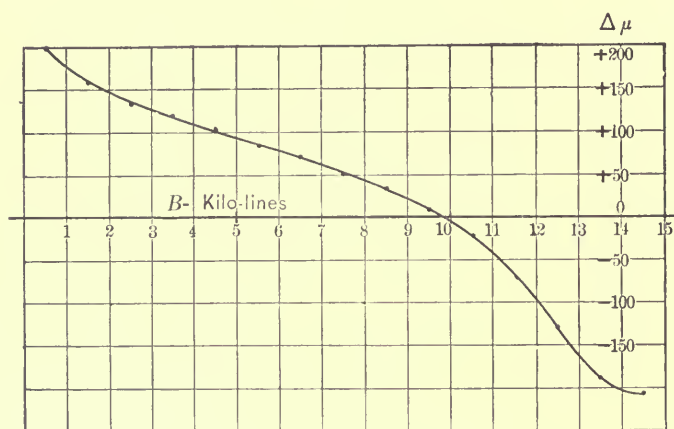


FIG. 58. First Differential Quotient of B, μ Curve

The e.m.f. acting on the circuit is $\frac{n}{x} e_0$, since $\frac{n}{x}$ cells of e.m.f. e_0 are in series. Therefore, the current delivered by the battery is,

$$i = \frac{\frac{n}{x} e_0}{\frac{nr_0}{x^2} + r},$$

and the power which this current produces in the resistance r , is,

$$P = ri^2 = \frac{rn^2 e_0^2}{x^2 \left(\frac{nr_0}{x^2} + r \right)^2}.$$

This is an extreme, if

$$y = \frac{nr_0}{x} + rx$$

is an extreme, hence,

$$\frac{dy}{dx} = -\frac{nr_0}{x^2} + r = 0,$$

and

$$x = \sqrt{\frac{nr_0}{r}} = 4;$$

that is, $x = \sqrt{\frac{nr_0}{r}} = 4$ cells are connected in multiple, and

$$\frac{n}{x} = \sqrt{\frac{nr}{r_0}} = 20 \text{ cells in series.}$$

115. Example 16, In an alternating-current transformer the loss of power is limited to 900 watts by the permissible temperature rise. The internal resistance of the transformer winding (primary, plus secondary reduced to the primary) is 2 ohms, and the core loss at 2000 volts impressed, is 400 watts, and varies with the 1.6th power of the magnetic density and therefore of the voltage. At what impressed voltage is the output of the transformer a maximum?

If e is the impressed e.m.f. and i is the current input, the power input into the transformer (approximately, at non-inductive load) is $P = ei$.

If the output is a maximum, at constant loss, the input P also is a maximum. The loss of power in the winding is $ri^2 = 2i^2$.

The loss of power in the iron at 2000 volts impressed is 400 watts, and at impressed voltage e it therefore is

$$\left(\frac{e}{2000}\right)^{1.6} \times 400,$$

and the total loss in the transformer, therefore, is

$$P_L = 2i^2 + 400\left(\frac{e}{2000}\right)^{1.6} = 900;$$

herefrom, it follows that,

$$i = \sqrt{450 - 200 \left(\frac{e}{2000} \right)^{1.6}},$$

and, substituting, into $P = ei$:

$$P = e \sqrt{450 - 200 \left(\frac{e}{2000} \right)^{1.6}}.$$

Simplified, this gives,

$$y = 2.25e^2 - \frac{e^{3.6}}{2000^{1.6}},$$

and, differentiating,

$$\frac{dy}{de} = 4.5e - \frac{3.6e^{2.6}}{2000^{1.6}} = 0,$$

and

$$\left(\frac{e}{2000} \right)^{1.6} = 1.25.$$

Hence,

$$\frac{e}{2000} = 1.15 \quad \text{and} \quad e = 2300 \text{ volts},$$

which, substituted, gives

$$P = 2300 \sqrt{450 - 200 \times 1.25} = 32.52 \text{ kw.}$$

116. Example 17. In a 5-kw. alternating-current transformer, at 1000 volts impressed, the core loss is 60 watts, the i^2r loss 150 watts. How must the impressed voltage be changed, to give maximum efficiency, (a) At full load of 5-kw; (b) at half load?

The core loss may be assumed as varying with the 1.6th power of the impressed voltage. If e is the impressed voltage, $i = \frac{5000}{e}$ is the current at full load, and $i_1 = \frac{2500}{e}$ is the current at half load, then at 1000 volts impressed, the full-load current is $\frac{5000}{e} = 5$ amperes, and since the i^2r loss is 150 watts, this gives

the internal resistance of the transformer as $r=6$ ohms, and herefrom the i^2r loss at impressed voltage e is respectively,

$$ri^2 = \frac{150 \times 10^6}{e^2} \quad \text{and} \quad ri_1^2 = \frac{37.5 \times 10^6}{e^2} \text{ watts.}$$

Since the core loss is 60 watts at 1000 volts, at the voltage e it is

$$P_i = 60 \times \left(\frac{e}{1000} \right)^{1.6} \text{ watts.}$$

The total loss, at full load, thus is

$$P_L = P_i + ri^2 = 60 \times \left(\frac{e}{1000} \right)^{1.6} + \frac{150 \times 10^6}{e^2},$$

and at half load it is

$$P_{L_1} = P_i + ri_1^2 = 60 \times \left(\frac{e}{1000} \right)^{1.6} + \frac{37.5 \times 10^6}{e^2}.$$

Simplified, this gives

$$y = \left(\frac{e}{1000} \right)^{1.6} + 2.5 \times 10^6 \cdot e^{-2},$$

$$y_1 = \left(\frac{e}{1000} \right)^{1.6} + 0.625 \times 10^6 \cdot e^{-2};$$

hence, differentiated,

$$1.6 \frac{e^{0.6}}{1000^{1.6}} - 5 \times 10^6 e^{-3} = 0;$$

$$1.6 \frac{e^{0.6}}{1000^{1.6}} - 1.25 \times 10^6 \cdot e^{-3} = 0;$$

$$e^{3.6} = 3.125 \times 10^6 \times 1000^{1.6} = 3.125 \times 10^{10.8};$$

$$e^{3.6} = 0.78125 \times 10^6 \times 1000^{1.6} = 0.78125 \times 10^{10.8};$$

hence, $e = 1373$ volts for maximum efficiency at full load.

and $e = 938$ volts for maximum efficiency at half load.

117. Example 18. (a) Constant voltage $e=1000$ is impressed upon a condenser of capacity $C=10$ mf., through a reactor of inductance $L=100$ nh., and a resistor of resistance $r=40$ ohms. What is the maximum value of the charging current?

(b) An additional resistor of resistance $r' = 210$ ohms is then inserted in series, making the total resistance of the condenser charging circuit, $r = 250$ ohms. What is the maximum value of the charging current?

The equation of the charging current of a condenser, through a circuit of low resistance, is ("Transient Electric Phenomena and Oscillations," p. 61):

$$i = \frac{2e}{q} \left\{ \epsilon^{-\frac{r}{2L}t} \sin \frac{q}{2L}t \right\},$$

where

$$q = \sqrt{\frac{4L}{C} - r^2},$$

and the equation of the charging current of a condenser, through a circuit of high resistance, is ("Transient Electric Phenomena and Oscillations," p. 51),

$$i = \frac{e}{s} \left\{ \epsilon^{-\frac{r-s}{2L}t} - \epsilon^{-\frac{r+s}{2L}t} \right\},$$

where

$$s = \sqrt{r^2 - \frac{4L}{C}}.$$

Substituting the numerical values gives:

$$(a) \quad i = 10.2 \epsilon^{-200t} \sin 980 t;$$

$$(b) \quad i = 6.667 \{ \epsilon^{-500t} - \epsilon^{-2000t} \}.$$

Simplified and differentiated, this gives:

$$(a) \quad z = \frac{di_1}{dt} = 4.9 \cos 980t - \sin 980t = 0;$$

hence

$$\tan 980t = 4.9$$

$$980t = 68.5^\circ = 1.20$$

$$t = \frac{1.20 + n\pi}{908} \text{ sec.}$$

$$(b) \quad z = \frac{di_2}{dt} = 4 \epsilon^{-2000t} - \epsilon^{-500t} = 0;$$

hence,

$$\epsilon + 1500t = 4,$$

$$1500t = \frac{\log 4}{\log \epsilon} = 1.38,$$

$$t = 0.00092 \text{ sec.},$$

and, by substituting these values of t into the equations of the current, gives the maximum values:

$$(a) \quad i = 10\epsilon^{-\frac{1.20 + n\pi}{4.9}} = 7.83\epsilon^{-0.64n} = 7.83 \times 0.53^n \text{ amperes};$$

that is, an infinite number of maxima, of gradually decreasing values: +7.83; -4.15; +2.20; -1.17 etc.

$$(b) \quad i = 6.667(\epsilon^{-0.46} - \epsilon^{-1.84}) = 3.16 \text{ amperes.}$$

118. Example 19. In an induction generator, the friction losses are $P_f = 100$ kw.; the iron loss is 200 kw. at the terminal voltage of $e = 4$ kv., and may be assumed as proportional to the 1.6th power of the voltage; the loss in the resistance of the conductors is 100 kw. at $i = 3000$ amperes output, and may be assumed as proportional to the square of the current, and the losses resulting from stray fields due to magnetic saturation are 100 kw. at $e = 4$ kv., and may in the range considered be assumed as approximately proportional to the 3.2th power of the voltage. Under what conditions of operation, regarding output, voltage and current, is the efficiency a maximum?

The losses may be summarized as follows:

$$\text{Friction loss,} \quad P_f = 100 \text{ kw.};$$

$$\text{Iron loss,} \quad P_i = 200 \left(\frac{e}{4} \right)^{1.6};$$

$$i^2 r \text{ loss,} \quad P_c = 100 \left(\frac{i}{3000} \right)^2;$$

$$\text{Saturation loss,} \quad P_s = 100 \left(\frac{e}{4} \right)^{3.2};$$

hence the total loss is $P_L = P_f + P_i + P_c + P_s$

$$= 100 \left\{ 1 + 2 \left(\frac{e}{4} \right)^{1.6} + \left(\frac{i}{3000} \right)^2 + \left(\frac{e}{4} \right)^{3.2} \right\}$$

The output is $P = ei$; hence, percentage of loss is

$$\lambda = \frac{P_L}{P} = \frac{100 \left\{ 1 + 2 \left(\frac{e}{4} \right)^{1.6} + \left(\frac{i}{3000} \right)^2 + \left(\frac{e}{4} \right)^{3.2} \right\}}{ei}.$$

The efficiency is a maximum, if the percentage loss λ is a minimum. For any value of the voltage e , this is the case at the current i , given by $\frac{d\lambda}{di} = 0$; hence, simplifying and differentiating λ ,

$$\frac{d\lambda}{di} = - \frac{1 + 2 \left(\frac{e}{4} \right)^{1.6} + \left(\frac{e}{4} \right)^{3.2}}{i^2} + \frac{1}{3000^2} = 0;$$

$$i = 3000 \sqrt{1 + 2 \left(\frac{e}{4} \right)^{1.6} + \left(\frac{e}{4} \right)^{3.2}};$$

then, substituting i in the expression of λ , gives

$$\lambda = \frac{1}{15e} \sqrt{1 + 2 \left(\frac{e}{4} \right)^{1.6} + \left(\frac{e}{4} \right)^{3.2}},$$

and λ is an extreme, if the simplified expression,

$$y = \frac{1}{e^2} + \frac{2}{4^{1.6} e^{0.4}} + \frac{1}{4^{3.2} e^{1.2}}$$

is an extreme, at

$$\frac{dy}{de} = -\frac{2}{e^3} - \frac{0.8}{4^{1.6} e^{1.4}} + \frac{1.2}{4^{3.2} e^{0.2}};$$

hence,

$$2 + \frac{0.8}{4^{1.6}} e^{1.6} - \frac{1.2}{4^{3.2}} e^{3.2} = 0;$$

hence,

$$\left(\frac{e}{4} \right)^{1.6} = \frac{2}{1.2} \quad \text{and} \quad e = 5.50 \text{ kv.},$$

and, by substitution the following values are obtained: $\lambda = 0.0323$; efficiency 96.77 per cent; current $i = 8000$ amperes; output $P = 44,000$ kw.

119. In all probability, this output is beyond the capacity of the generator, as limited by heating. The foremost limitation probably will be the $i^2 r$ heating of the conductors; that is,

the maximum permissible current will be restricted to, for instance, $i = 5000$ amperes.

For any given value of current i , the maximum efficiency, that is, minimum loss, is found by differentiating,

$$\lambda = \frac{100 \left\{ 1 + 2 \left(\frac{e}{4} \right)^{1.6} + \left(\frac{i}{3000} \right)^2 + \left(\frac{e}{4} \right)^{3.2} \right\}}{ei}$$

over e , thus:

$$\frac{d\lambda}{de} = 0.$$

Simplified, λ gives

$$y = \frac{1}{e} \left\{ 1 + \left(\frac{i}{3000} \right)^2 \right\} + \frac{2}{4^{1.6}} e^{0.6} + \frac{1}{4^{3.2}} e^{2.2};$$

hence, differentiated, it gives

$$\frac{dy}{de} = -\frac{1}{e^2} \left\{ 1 + \left(\frac{i}{3000} \right)^2 \right\} + \frac{1.2}{4^{1.6}} e^{-0.4} + \frac{2.2 e^{1.2}}{4^{3.2}} = 0;$$

$$\left(\frac{e}{4} \right)^{3.2} + \frac{6}{11} \left(\frac{e}{4} \right)^{1.6} = \frac{5}{11} \left\{ 1 + \left(\frac{i}{3000} \right)^2 \right\};$$

$$\left(\frac{e}{4} \right)^{1.6} = \frac{-3 + \sqrt{64 + 55 \frac{i^2}{3000^2}}}{11};$$

For $i = 5000$, this gives:

$$\left(\frac{e}{4} \right)^{1.6} = 1.065 \quad \text{and} \quad e = 4.16 \text{ kv.};$$

hence,

$\lambda = 0.0338$, Efficiency 96.62 per cent, Power $P = 20,800$ kw.

Method of Least Squares.

120. An interesting and very important application of the theory of extremes is given by the *method of least squares*, which is used to calculate the most accurate values of the constants of functions from numerical observations which are more numerous than the constants.

If $y = f(x)$, (1)

is a function having the constants $a, b, c \dots$ and the form of the function (1) is known, for instance,

$$y = a + bx + cx^2, \quad . \quad . \quad . \quad . \quad . \quad . \quad (2)$$

and the constants a, b, c are not known, but the numerical values of a number of corresponding values of x and y are given, for instance, by experiment, $x_1, x_2, x_3, x_4 \dots$ and $y_1, y_2, y_3, y_4 \dots$, then from these corresponding numerical values x_n and y_n the constants $a, b, c \dots$ can be calculated, if the numerical values, that is, the observed points of the curve, are sufficiently numerous.

If less points $x_1, y_1, x_2, y_2 \dots$ are observed, then the equation (1) has constants, obviously these constants cannot be calculated, as not sufficient data are available therefor.

If the number of observed points equals the number of constants, they are just sufficient to calculate the constants. For instance, in equation (2), if three corresponding values $x_1, y_1; x_2, y_2; x_3, y_3$ are observed, by substituting these into equation (2), three equations are obtained:

$$\left. \begin{aligned} y_1 &= a + bx_1 + cx_1^2; \\ y_2 &= a + bx_2 + cx_2^2; \\ y_3 &= a + bx_3 + cx_3^2, \end{aligned} \right\} \quad . \quad . \quad . \quad . \quad . \quad . \quad (3)$$

which are just sufficient for the calculation of the three constants a, b, c .

Three observations would therefore be sufficient for determining three constants, if the observations were absolutely correct. This, however, is not the case, but the observations always contain *errors of observation*, that is, unavoidable inaccuracies, and constants calculated by using only as many observations as there are constants, are not very accurate.

Thus, in experimental work, always more observations are made than just necessary for the determination of the constants, for the purpose of getting a higher accuracy. Thus, for instance, in astronomy, for the calculation of the orbit of a comet, less than four observations are theoretically sufficient, but if possible hundreds are taken, to get a greater accuracy in the determination of the constants of the orbit.

If, then, for the determination of the constants a, b, c of equation (2), six pairs of corresponding values of x and y were determined, any three of these pairs would be sufficient to give a, b, c , as seen above, but using different sets of three observations, would not give the same values of a, b, c (as it should, if the observations were absolutely accurate), but different values, and none of these values would have as high an accuracy as can be reached from the experimental data, since none of the values uses all observations.

$$121. \text{ If } y=f(x), \quad . \quad . \quad . \quad . \quad . \quad . \quad (1)$$

is a function containing the constants $a, b, c \dots$, which are still unknown, and $x_1, y_1; x_2, y_2; x_3, y_3; \text{ etc.}$, are corresponding experimental values, then, if these values were absolutely correct, and the correct values of the constants $a, b, c \dots$ chosen, $y_1=f(x_1)$ would be true; that is,

$$\left. \begin{array}{l} f(x_1) - y_1 = 0; \\ f(x_2) - y_2 = 0, \text{ etc.} \end{array} \right\} . \quad . \quad . \quad . \quad . \quad . \quad (5)$$

Due to the errors of observation, this is not the case, but even if $a, b, c \dots$ are the correct values,

$$y_1 \neq f(x_1) \text{ etc.}; \quad . \quad . \quad . \quad . \quad . \quad . \quad (6)$$

that is, a small difference, or error, exists, thus

$$\left. \begin{array}{l} f(x_1) - y_1 = \delta_1; \\ f(x_2) - y_2 = \delta_2, \text{ etc.}; \end{array} \right\} . \quad . \quad . \quad . \quad . \quad . \quad (7)$$

If instead of the correct values of the constants, $a, b, c \dots$, other values were chosen, different errors $\delta_1, \delta_2 \dots$ would obviously result.

From probability calculation it follows, that, if the correct values of the constants $a, b, c \dots$ are chosen, the sum of the squares of the errors,

$$\delta_1^2 + \delta_2^2 + \delta_3^2 + \dots \quad . \quad . \quad . \quad . \quad . \quad . \quad (8)$$

is less than for any other value of the constants $a, b, c \dots$; that is, it is a minimum.

122. The problem of determining the constants $a, b, c \dots$, thus consists in finding a set of constants, which makes the sum of the squares of the errors δ a minimum; that is,

$$z = \Sigma \delta^2 = \text{minimum}, \dots \dots \dots (9)$$

is the requirement, which gives the most accurate or most probable set of values of the constants $a, b, c \dots$.

Since by (7), $\delta = f(x) - y$, it follows from (9) as the condition, which gives the most probable value of the constants $a, b, c \dots$:

$$z = \Sigma \{f(x) - y\}^2 = \text{minimum}; \dots \dots \dots (10)$$

that is, the least sum of the squares of the errors gives the most probable value of the constants $a, b, c \dots$.

To find the values of $a, b, c \dots$, which fulfill equation (10), the differential quotients of (10) are equated to zero, and give

$$\left. \begin{aligned} \frac{dz}{da} &= \Sigma \{f(x) - y\} \frac{df(x)}{da} = 0; \\ \frac{dz}{db} &= \Sigma \{f(x) - y\} \frac{df(x)}{db} = 0; \\ \frac{dz}{dc} &= \Sigma \{f(x) - y\} \frac{df(x)}{dc} = 0; \text{ etc.} \end{aligned} \right\} \dots \dots \dots (11)$$

This gives as many equations as there are constants $a, b, c \dots$, and therefore just suffices for their calculation, and the values so calculated are the most probable, that is, the most accurate values.

Where extremely high accuracy is required, as for instance in astronomy when calculating from observations extending over a few months only, the orbit of a comet which possibly lasts thousands of years, the method of least squares must be used, and is frequently necessary also in engineering, to get from a limited number of observations the highest accuracy of the constants.

123. As instance, the method of least squares may be applied in separating from the observations of an induction motor, when running light, the component losses, as friction, hysteresis, etc.

In a 440-volt 50-h.p. induction motor, when running light, that is, without load, at various voltages, let the terminal voltage e , the current input i , and the power input p be observed as given in the first three columns of Table I:

TABLE I

e	i	p	i^2r	p_0	$\frac{p_0}{\text{calc.}}$	J
148	8	790	13	780	746	+ 32
220	11	920	24	900	962	- 62
320	19	1500	72	1430	1382	+ 48
410	23	1920	106	1810	1875	- 35
440	26	2220	135	2085	2058	+ 27
473	29	2450	168	2280	2280	0
580	43	3700	370	3330	3080	+ 250
640	56	5000	627	4370	3600	+ 770
700	75	8000	1125	6875	4150	+ 2725

The power consumed by the motor while running light consists of: The friction loss, which can be assumed as constant, a ; the hysteresis loss, which is proportional to the 1.6th power of the magnetic flux, and therefore of the voltage, $be^{1.6}$; the eddy current losses, which are proportional to the square of the magnetic flux, and therefore of the voltage, ce^2 ; and the i^2r loss in the windings. The total power is,

$$p = a + be^{1.6} + ce^2 + ri^2. \quad (12)$$

From the resistance of the motor windings, $r = 0.2$ ohm, and the observed values of current i , the i^2r loss is calculated, and tabulated in the fourth column of Table I, and subtracted from p , leaving as the total mechanical and magnetic losses the values of p_0 given in the fifth column of the table, which should be expressed by the equation:

$$p_0 = a + be^{1.6} + ce^2. \quad (13)$$

This leaves three constants, a , b , c , to be calculated.

Plotting now in Fig. 59 with values of e as abscissas, the current i and the power p_0 give curves, which show that within the voltage range of the test, a change occurs in the motor,

as indicated by the abrupt rise of current and of power beyond 473 volts. This obviously is due to beginning magnetic saturation of the iron structure. Since with beginning saturation a change of the magnetic distribution must be expected, that is, an increase of the magnetic stray field and thereby increase of eddy current losses, it is probable that at this point the con-

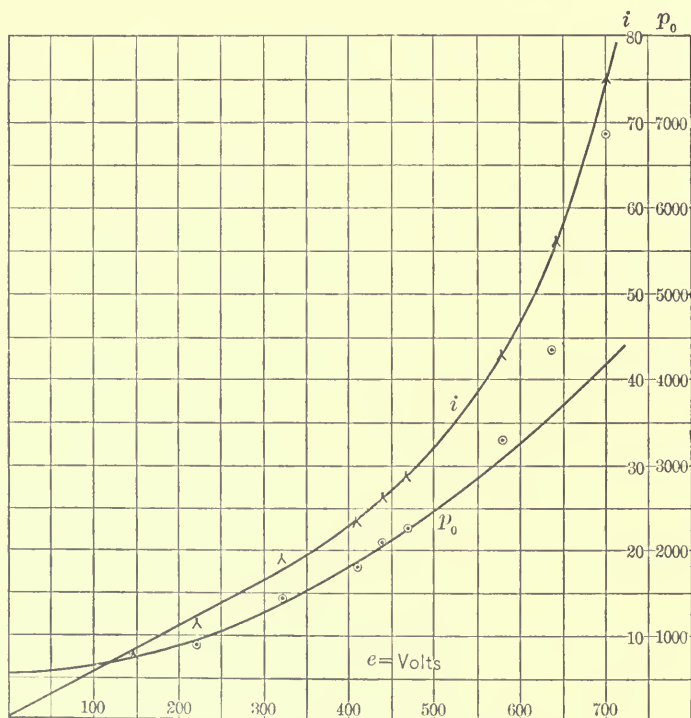


FIG. 59. Excitation Power of Induction Motor.

stants in equation (13) change, and no set of constants can be expected to represent the entire range of observation. For the calculation of the constants in (13), thus only the observations below the range of magnetic saturation can safely be used, that is, up to 473 volts.

From equation (13) follows as the error of an individual observation of e and p_0 :

$$\delta = a + be^{1.6} + ce^2 - p_0; \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (14)$$

hence,

$$z = \Sigma \delta^2 = \Sigma \{a + be^{1.6} + ce^2 - p_0\}^2 = \text{minimum}, \quad (15)$$

thus:

$$\left. \begin{aligned} \frac{dz}{da} &= \Sigma \{a + be^{1.6} + ce^2 - p_0\} = 0; \\ \frac{dz}{db} &= \Sigma \{a + be^{1.6} + ce^2 - p_0\} e^{1.6} = 0; \\ \frac{dz}{dc} &= \Sigma \{a + be^{1.6} + ce^2 - p_0\} e^2 = 0; \end{aligned} \right\} \quad . \quad . \quad . \quad (16)$$

and, if n is the number of observations used ($n=6$ in this instance, from $e=148$ to $e=473$), this gives the following equations:

$$\left. \begin{aligned} na + b \Sigma e^{1.6} + c \Sigma e^2 - \Sigma p_0 &= 0; \\ a \Sigma e^{1.6} + b \Sigma e^{3.2} + c \Sigma e^{3.6} - \Sigma e^{1.6} p_0 &= 0; \\ a \Sigma e^2 + b \Sigma e^{3.6} + c \Sigma e^4 - \Sigma e^2 p_0 &= 0. \end{aligned} \right\} \quad . \quad . \quad . \quad (17)$$

Substituting in (17) the numerical values from Table I gives,

$$\left. \begin{aligned} a + 11.7 b 10^3 + 126 c 10^3 &= 1550; \\ a + 11.6 b 10^3 + 163 c 10^3 &= 1830; \\ a + 15.1 b 10^3 + 170 c 10^3 &= 1880; \end{aligned} \right\} \quad . \quad . \quad . \quad (18)$$

hence,

$$\left. \begin{aligned} a &= 540; \\ b &= 32.5 \times 10^{-3}; \\ c &= 5 \times 10^{-3}, \end{aligned} \right\} \quad . \quad . \quad . \quad . \quad . \quad . \quad (19)$$

and

$$p_0 = 540 + 0.0325 e^{1.6} + 0.005 e^2, \quad . \quad . \quad . \quad . \quad (20)$$

The values of p_0 , calculated from equation (20), are given in the sixth column of Table I, and their differences from the observed values in the last column. As seen, the errors are in both directions from the calculated values, except for the three highest voltages, in which the observed values rapidly increase beyond the calculated, due probably to the appearance of a

loss which does not exist at lower voltages—the eddy currents caused by the magnetic stray field of saturation.

This rapid divergency of the observed from the calculated values at high voltages shows that a calculation of the constants, based on all observations, would have led to wrong values, and demonstrates the necessity, first, to critically review the series of observations, before using them for deriving constants, so as to exclude constant errors or unidirectional deviation. It must be realized that the method of least squares gives the most probable value, that is, the most accurate results derivable from a series of observations, only so far as the accidental errors of observations are concerned, that is, such errors which follow the general law of probability. The method of least squares, however, cannot eliminate constant errors, that is, deviation of the observations which have the tendency to be in one direction, as caused, for instance, by an instrument reading too high, or too low, or the appearance of a new phenomenon in a part of the observation, as an additional loss in above instance, etc. Against such constant errors only a critical review and study of the method and the means of observation can guard, that is, judgment, and not mathematical formalism.

The method of least squares gives the highest accuracy available with a given number of observations, but is frequently very laborious, especially if a number of constants are to be calculated. It, therefore, is mainly employed where the number of observations is limited and cannot be increased at will; but where it can be increased by taking some more observations—as is generally the case with experimental engineering investigations—the same accuracy is usually reached in a shorter time by taking a few more observations and using a simpler method of calculation of the constants, as the $\Sigma\Delta$ -method described in paragraphs 153 to 157.

Diophantic Equations.

123A.—The method of least squares deals with the case, when there are more equations than unknown quantities. In this case, there exists no set of values of the unknown quantities, which exactly satisfies the equations, and the problem is, to find

the set of values, which most nearly satisfies the equations, that is, which is the most probable.

Inversely, sometimes in engineering the case is met, when there are more unknown than equations, for instance, two equations with three unknown quantities. Mathematically, this gives not one, but an infinite series of sets of solutions of the equations. Physically however in such a case, the number of permissible solutions may be limited by some condition outside of the algebra of equations. Such for instance often is, in physics, engineering, etc., the condition that the values of the unknown quantities must be positive integer numbers.

Thus an engineering problem may lead to two equations with three unknown quantities, which latter are limited by the condition of being positive and integer, or similar requirements, and in such a case, the number of solutions of the equation may be finite, although there are more equations than unknown quantities.

For instance:

In calculating from economic consideration, in a proposed hydroelectric generating station, the number of generators, exciters and step-up transformers, let:

x = number of generators

y = number of exciters

z = number of transformers

Suppose now, the physical and economic conditions of the installation lead us to the equations:

$$8x + 3y + z = 49 \quad (1)$$

$$2x + y + 3z = 21 \quad (2)$$

These are two equations with three unknown, x , y , z ; these unknown however are conditioned by the physical requirement, that they are integer positive numbers.

To attempt to secure a third equation would then over determine the problem, and give either wrong, or limited results.

Eliminating z from (1) and (2), gives:

$$11x + 4y = 63 \quad (3)$$

hence:

$$y = \frac{63 - 11x}{4} = 15 - 2x + \frac{3 - 3x}{4} \quad (4)$$

since y must be an integer number, $\frac{3 - 3x}{4}$ must also be an integer number. Call this u , it is:

$$\begin{aligned} \frac{3 - 3x}{4} &= u \\ x &= \frac{3 - 4u}{3} = 1 - u - \frac{u}{3} \end{aligned} \quad (5)$$

since x must be an integer number, $\frac{u}{3}$ must also be an integer number, that is:

$$u = 3v \quad (6)$$

hence, substituted into (5), (4) and (2):

$$\left. \begin{aligned} x &= 1 - 4v \\ y &= 13 + 11v \\ z &= 2 - v \end{aligned} \right\} \quad (7)$$

(7) thus are the solutions of the equations (1) (2), where v is any integer number.

As seen, mathematically, there are an infinite number of solutions.

Substituting now for v integer numbers:

$v = +2$	$+1$	0	-1	-2
$x = -7$	-3	$+1$	$+5$	$+9$
$y = +35$	$+24$	$+13$	$+2$	-9
$z = 0$	$+1$	$+2$	$+3$	$+4$

As seen, there are only two solutions, for $v = 0$ and $v = -1$, which give for x , y , and z , three integer positive values, and which thus satisfy the physical restriction.

$v = 0$; $x = 1$, $y = 13$, $z = 2$ is excluded by engineering consideration, as nobody would consider thirteen exciters with one generator, and thus there remains only one applicable solution:

$$x = 5$$

$$y = 2$$

$$z = 3$$

We thus have here the case of two equations with three unknown quantities, which have only one single set of these unknown quantities satisfying the problem, and thus give a definite solution, though mathematically indefinite.

This type of equation has first been studied by Diophantes of Alexandria.



CHAPTER V.

METHODS OF APPROXIMATION.

124. The investigation even of apparently simple engineering problems frequently leads to expressions which are so complicated as to make the numerical calculations of a series of values very cumbersome and almost impossible in practical work. Fortunately in many such cases of engineering problems, and especially in the field of electrical engineering, the different quantities which enter into the problem are of very different magnitude. Many apparently complicated expressions can frequently be greatly simplified, to such an extent as to permit a quick calculation of numerical values, by neglecting terms which are so small that their omission has no appreciable effect on the accuracy of the result; that is, leaves the result correct within the limits of accuracy required in engineering, which usually, depending on the nature of the problem, is not greater than from 0.1 per cent to 1 per cent.

Thus, for instance, the voltage consumed by the resistance of an alternating-current transformer is at full load current only a small fraction of the supply voltage, and the exciting current of the transformer is only a small fraction of the full load current, and, therefore, the voltage consumed by the exciting current in the resistance of the transformer is only a small fraction of a small fraction of the supply voltage, hence, it is negligible in most cases, and the transformer equations are greatly simplified by omitting it. The power loss in a large generator or motor is a small fraction of the input or output, the drop of speed at load in an induction motor or direct-current shunt motor is a small fraction of the speed, etc., and the square of this fraction can in most cases be neglected, and the expression simplified thereby.

Frequently, therefore, in engineering expressions containing small quantities, the products, squares and higher

powers of such quantities may be dropped and the expression thereby simplified; or, if the quantities are not quite as small as to permit the neglect of their squares, or where a high accuracy is required, the first and second powers may be retained and only the cubes and higher powers dropped.

The most common method of procedure is, to resolve the expression into an infinite series of successive powers of the small quantity, and then retain of this series only the first term, or only the first two or three terms, etc., depending on the smallness of the quantity and the required accuracy.

125. The forms most frequently used in the reduction of expressions containing small quantities are multiplication and division, the binomial series, the exponential and the logarithmic series, the sine and the cosine series, etc.

Denoting a small quantity by s , and where several occur, by $s_1, s_2, s_3 \dots$ the following expression holds:

$$(1 \pm s_1)(1 \pm s_2) = 1 \pm s_1 \pm s_2 \pm s_1 s_2,$$

and, since $s_1 s_2$ is small compared with the small quantities s_1 and s_2 , or, as usually expressed, $s_1 s_2$ is a small quantity of higher order (in this case of second order), it may be neglected, and the expression written:

$$(1 \pm s_1)(1 \pm s_2) = 1 \pm s_1 \pm s_2. \quad \dots \quad (1)$$

This is one of the most useful simplifications: the multiplication of terms containing small quantities is replaced by the simple addition of the small quantities.

If the small quantities s_1 and s_2 are not added (or subtracted) to 1, but to other finite, that is, not small quantities a and b , a and b can be taken out as factors, thus,

$$(a \pm s_1)(b \pm s_2) = ab \left(1 \pm \frac{s_1}{a}\right) \left(1 \pm \frac{s_2}{b}\right) = ab \left(1 \pm \frac{s_1}{a} \pm \frac{s_2}{b}\right), \quad (2)$$

where $\frac{s_1}{a}$ and $\frac{s_2}{b}$ must be small quantities.

As seen, in this case, s_1 and s_2 need not necessarily be absolutely small quantities, but may be quite large, provided that a and b are still larger in magnitude; that is, s_1 must be small compared with a , and s_2 small compared with b . For instance,

in astronomical calculations the mass of the earth (which absolutely can certainly not be considered a small quantity) is neglected as small quantity compared with the mass of the sun. Also in the effect of a lightning stroke on a primary distribution circuit, the normal line voltage of 2200 may be neglected as small compared with the voltage impressed by lightning, etc.

126. Example. In a direct-current shunt motor, the impressed voltage is $e_0=125$ volts; the armature resistance is $r_0=0.02$ ohm; the field resistance is $r_1=50$ ohms; the power consumed by friction is $p_f=300$ watts, and the power consumed by iron loss is $p_i=400$ watts. What is the power output of the motor at $i_0=50, 100$ and 150 amperes input?

The power produced at the armature conductors is the product of the voltage e generated in the armature conductors, and the current i through the armature, and the power output at the motor pulley is,

$$p = ei - p_f - p_i. \quad (3)$$

The current in the motor field is $\frac{e_0}{r_1}$, and the armature current therefore is,

$$i = i_0 - \frac{e_0}{r_1}, \quad (4)$$

where $\frac{e_0}{r_1}$ is a small quantity, compared with i_0 .

The voltage consumed by the armature resistance is $r_0 i$, and the voltage generated in the motor armature thus is:

$$e = e_0 - r_0 i, \quad (5)$$

where $r_0 i$ is a small quantity compared with e_0 .

Substituting herein for i the value (4) gives,

$$e = e_0 - r_0 \left(i_0 - \frac{e_0}{r_1} \right). \quad (6)$$

Since the second term of (6) is small compared with e_0 , and in this second term, the second term $\frac{e_0}{r_1}$ is small compared with i_0 , it can be neglected as a small term of higher

order; that is, as small compared with a small term, and expression (6) simplified to

$$e = e_0 - r_0 i_0. \quad . \quad . \quad . \quad . \quad . \quad . \quad (7)$$

Substituting (4) and (7) into (3) gives,

$$\begin{aligned} p &= (e_0 - r_0 i_0) \left(i_0 - \frac{e_0}{r_1} \right) - p_f - p_i \\ &= e_0 i_0 \left(1 - \frac{r_0 i_0}{e_0} \right) \left(1 - \frac{e_0}{r_1 i_0} \right) - p_f - p_i. \quad . \quad . \quad . \quad (8) \end{aligned}$$

Expression (8) contains a product of two terms with small quantities, which can be multiplied by equation (1), and thereby gives,

$$\begin{aligned} p &= e_0 i_0 \left(1 - \frac{r_0 i_0}{e_0} - \frac{e_0}{r_1 i_0} \right) - p_f - p_i \\ &= e_0 i_0 - r_0 i_0^2 - \frac{e_0^2}{r_1} - p_f - p_i. \quad . \quad . \quad . \quad . \quad . \quad (9) \end{aligned}$$

Substituting the numerical values gives,

$$\begin{aligned} p &= 125 i_0 - 0.02 i_0^2 - 562.5 - 300 - 400 \\ &= 125 i_0 - 0.02 i_0^2 - 1260 \text{ approximately;} \end{aligned}$$

thus, for $i_0 = 50, 100$, and 150 amperes; $p = 4940, 11,040$, and $17,040$ watts respectively.

127. Expressions containing a small quantity in the denominator are frequently simplified by bringing the small quantity in the numerator, by division as discussed in Chapter II paragraph 39, that is, by the series,

$$\frac{1}{1 \pm x} = 1 \mp x + x^2 \mp x^3 + x^4 \mp x^5 + \dots \quad (10)$$

which series, if x is a small quantity s , can be approximated by:

$$\left. \begin{aligned} \frac{1}{1+s} &= 1-s; \\ \frac{1}{1-s} &= 1+s; \end{aligned} \right\}, \quad . \quad . \quad . \quad . \quad . \quad (11)$$

or, where a greater accuracy is required,

$$\left. \begin{aligned} \frac{1}{1+s} &= 1 - s + s^2; \\ \frac{1}{1-s} &= 1 + s + s^2. \end{aligned} \right\} \dots \dots \dots (12)$$

By the same expressions (11) and (12) a small quantity contained in the numerator may be brought into the denominator where this is more convenient, thus:

$$\left. \begin{aligned} 1+s &= \frac{1}{1-s}; \\ 1-s &= \frac{1}{1+s}; \text{ etc.} \end{aligned} \right\} \dots \dots \dots (13)$$

More generally then, an expression like $\frac{b}{a \pm s}$, where s is small compared with a , may be simplified by approximation to the form,

$$\frac{b}{a \pm s} = \frac{b}{a \left(1 \pm \frac{s}{a} \right)} = \frac{b}{a} \left(1 \mp \frac{s}{a} \right); \dots \dots \dots (14)$$

or, where a greater exactness is required, by taking in the second term,

$$\frac{b}{a \pm s} = \frac{b}{a} \left(1 \mp \frac{s}{a} + \frac{s^2}{a^2} \right) \dots \dots \dots (15)$$

128. Example. What is the current input to an induction motor, at impressed voltage e_0 and slip s (given as fraction of synchronous speed) if $r_0 + jx_0$ is the impedance of the primary circuit of the motor, and $r_1 + jx_1$ the impedance of the secondary circuit of the motor at full frequency, and the exciting current of the motor is neglected; assuming s to be a small quantity; that is, the motor running at full speed?

Let E be the e.m.f. generated by the mutual magnetic flux, that is, the magnetic flux which interlinks with primary and with secondary circuit, in the primary circuit. Since the frequency of the secondary circuit is the fraction s of the frequency

of the primary circuit, the generated e.m.f. of the secondary circuit is $s\dot{E}$.

Since x_1 is the reactance of the secondary circuit at full frequency, at the fraction s of full frequency the reactance of the secondary circuit is sx_1 , and the impedance of the secondary circuit at slip s , therefore, is $r_1 + jsx_1$; hence the secondary current is,

$$\dot{I} = \frac{s\dot{E}}{r_1 + jsx_1}.$$

If the exciting current is neglected, the primary current equals the secondary current (assuming the secondary of the same number of turns as the primary, or reduced to the same number of turns); hence, the current input into the motor is

$$\dot{I} = \frac{s\dot{E}}{r_1 + jsx_1} \dots \dots \dots (16)$$

The second term in the denominator is small compared with the first term, and the expression (16) thus can be approximated by

$$\dot{I} = \frac{s\dot{E}}{r_1 \left(1 + j \frac{sx_1}{r_1}\right)} = \frac{s\dot{E}}{r_1} \left(1 - j \frac{sx_1}{r_1}\right) \dots \dots \dots (17)$$

The voltage $\dot{E}$ generated in the primary circuit equals the impressed voltage e_0 , minus the voltage consumed by the current $\dot{I}$ in the primary impedance; $r_0 + jx_0$ thus is

$$\dot{E} = e_0 - \dot{I}(r_0 + jx_0) \dots \dots \dots (18)$$

Substituting (17) into (18) gives

$$\dot{E} = e_0 - \frac{s\dot{E}}{r_1}(r_0 + jx_0) \left(1 - j \frac{sx_1}{r_1}\right) \dots \dots \dots (19)$$

In expression (19), the second term on the right-hand side, which is the impedance drop in the primary circuit, is small compared with the first term e_0 , and in the factor $\left(1 - j \frac{sx_1}{r_1}\right)$ of this small term, the small term $j \frac{sx_1}{r_1}$ can thus be neglected

as a small term of higher order, and equation (19) abbreviated to

$$\dot{E} = e_0 - \frac{s\dot{E}}{r_1}(r_0 + jx_0). \quad (20)$$

From (20) it follows that

$$\dot{E} = \frac{e_0}{1 + \frac{s}{r_1}(r_0 + jx_0)},$$

and by (13),

$$\dot{E} = e_0 \left\{ 1 - \frac{s}{r_1}(r_0 + jx_0) \right\}. \quad (21)$$

Substituting (21) into (17) gives

$$I = \frac{se_0}{r_1} \left(1 - j\frac{sX_1}{r_1} \right) \left\{ 1 - \frac{s}{r_1}(r_0 + jx_0) \right\},$$

and by (1),

$$\begin{aligned} I &= \frac{se_0}{r_1} \left\{ 1 - j\frac{sX_1}{r_1} - \frac{s}{r_1}(r_0 + jx_0) \right\} \\ &= \frac{se_0}{r_1} \left\{ 1 - s\frac{r_0}{r_1} - js\frac{x_0 + jx_1}{r_1} \right\}. \end{aligned} \quad (22)$$

If then, $I_{00} = i_0 - ji_0'$ is the exciting current, the total current input into the motor is, approximately,

$$\begin{aligned} I_0 &= I + I_{00} \\ &= \frac{se_0}{r_1} \left\{ 1 + s\frac{r_0}{r_1} - js\frac{x_0 + jx_1}{r_1} \right\} + i_0 - ji_0'. \end{aligned} \quad (23)$$

129. One of the most important expressions used for the reduction of small terms is the binomial series:

$$\begin{aligned} (1 \pm x)^n &= 1 \pm nx + \frac{n(n-1)}{2}x^2 \pm \frac{n(n-1)(n-2)}{3}x^3 \\ &\quad + \frac{n(n-1)(n-2)(n-3)}{4}x^4 \pm \dots \end{aligned} \quad (24)$$

If x is a small term s , this gives the approximation,

$$(1 \pm s)^n = 1 \pm ns; \quad (25)$$

or, using the second term also, it gives

$$(1 \pm s)^n = 1 \pm ns + \frac{n(n-1)}{2} s^2. \quad \dots \quad (26)$$

In a more general form, this expression gives

$$(a \pm s)^n = a^n \left(1 \pm \frac{s}{a}\right)^n = a^n \left(1 \pm \frac{ns}{a}\right); \text{ etc.} \quad \dots \quad (27)$$

By the binomial, higher powers of terms containing small quantities, and, assuming n as a fraction, roots containing small quantities, can be eliminated; for instance,

$$\sqrt[n]{a \pm s} = (a \pm s)^{\frac{1}{n}} = a^{\frac{1}{n}} \left(1 \pm \frac{s}{a}\right)^{\frac{1}{n}} = \sqrt[n]{a} \left(1 \pm \frac{s}{na}\right);$$

$$\frac{1}{(a \pm s)^n} = \frac{1}{a^n \left(1 \pm \frac{s}{a}\right)^n} = \frac{1}{a^n} \left(1 \pm \frac{s}{a}\right)^{-n} = \frac{1}{a^n} \left(1 \mp \frac{ns}{a}\right);$$

$$\frac{1}{\sqrt[n]{(a \pm s)}} = (a \pm s)^{-\frac{1}{n}} = a^{-\frac{1}{n}} \left(1 \pm \frac{s}{a}\right)^{-\frac{1}{n}} = \frac{1}{\sqrt[n]{a}} \left(1 \mp \frac{s}{na}\right);$$

$$\sqrt[n]{(a \pm s)^m} = (a \pm s)^{\frac{m}{n}} = a^{\frac{m}{n}} \left(1 \pm \frac{s}{a}\right)^{\frac{m}{n}} = \sqrt[n]{a^m} \left(1 \pm \frac{ms}{na}\right); \text{ etc.}$$

One of the most common uses of the binomial series is for the elimination of squares and square roots, and very frequently it can be conveniently applied in mere numerical calculations; as, for instance,

$$(201)^2 = 200^2 \left(1 + \frac{1}{200}\right)^2 = 40,000 \left(1 + \frac{1}{100}\right) = 40,400;$$

$$29.9^2 = 30^2 \left(1 - \frac{1}{300}\right)^2 = 900 \left(1 - \frac{1}{150}\right) = 900 - 6 = 894;$$

$$\sqrt{99.8} = 10\sqrt{1 - 0.02} = 10(1 - 0.02)^{\frac{1}{2}} = 10(1 - 0.01) = 9.99;$$

$$\frac{1}{\sqrt{1.03}} = \frac{1}{(1 + 0.03)^{\frac{1}{2}}} = \frac{1}{1.015} = 0.985; \text{ etc.}$$

130. Example 1. If r is the resistance, x the reactance of an alternating-current circuit with impressed voltage e , the current is

$$i = \frac{e}{\sqrt{r^2 + x^2}}.$$

If the reactance x is small compared with the resistance r , as is the case in an incandescent lamp circuit, then,

$$\begin{aligned} i &= \frac{e}{\sqrt{r^2 + x^2}} = \frac{e}{r\sqrt{1 + \left(\frac{x}{r}\right)^2}} = \frac{e}{r} \left\{ 1 + \left(\frac{x}{r}\right)^2 \right\}^{-\frac{1}{2}} \\ &= \frac{e}{r} \left\{ 1 - \frac{1}{2} \left(\frac{x}{r}\right)^2 \right\}. \end{aligned}$$

If the resistance is small compared with the reactance, as is the case in a reactive coil, then,

$$\begin{aligned} i &= \frac{e}{\sqrt{r^2 + x^2}} = \frac{e}{x\sqrt{1 + \left(\frac{r}{x}\right)^2}} = \frac{e}{x} \left\{ 1 + \left(\frac{r}{x}\right)^2 \right\}^{-\frac{1}{2}} \\ &= \frac{e}{x} \left\{ 1 - \frac{1}{2} \left(\frac{r}{x}\right)^2 \right\}. \quad \dots \dots \dots (28) \end{aligned}$$

Example 2. How does the short-circuit current of an alternator vary with the speed, at constant field excitation?

When an alternator is short circuited, the total voltage generated in its armature is consumed by the resistance and the synchronous reactance of the armature.

The voltage generated in the armature at constant field excitation is proportional to its speed. Therefore, if e_0 is the voltage generated in the armature at some given speed S_0 , for instance, the rated speed of the machine, the voltage generated at any other speed S is

$$e = \frac{S}{S_0} e_0;$$

or, if for convenience, the fraction $\frac{S}{S_0}$ is denoted by a , then

$$a = \frac{S}{S_0} \quad \text{and} \quad e = ae_0,$$

where a is the ratio of the actual speed, to that speed at which the generated voltage is e_0 .

If r is the resistance of the alternator armature, x_0 the synchronous reactance at speed S_0 , the synchronous reactance at speed S is $x = ax_0$, and the current at short circuit then is

$$i = \frac{e}{\sqrt{r^2 + x^2}} = \frac{ae_0}{\sqrt{r^2 + a^2x_0^2}} \quad \dots \quad (29)$$

Usually r and x_0 are of such magnitude that r consumes at full load about 1 per cent or less of the generated voltage, while the reactance voltage of x_0 is of the magnitude of from 20 to 50 per cent. Thus r is small compared with x_0 , and if a is not very small, equation (29) can be approximated by

$$i = \frac{ae_0}{ax_0\sqrt{1 + \left(\frac{r}{ax_0}\right)^2}} = \frac{e_0}{x_0} \left\{ 1 - \frac{1}{2} \left(\frac{r}{ax_0} \right)^2 \right\} \quad \dots \quad (30)$$

Then if $x_0 = 20r$, the following relations exist:

$a =$	0.2	0.5	1.0	2.0
$i = \frac{e_0}{x_0} \times$	0.9688	0.995	0.99875	0.99969

That is, the short-circuit current of an alternator is practically constant independent of the speed, and begins to decrease only at very low speeds.

131. Exponential functions, logarithms, and trigonometric functions are the ones frequently met in electrical engineering.

The exponential function is defined by the series,

$$e^{\pm x} = 1 \pm x + \frac{x^2}{2} \pm \frac{x^3}{3} + \frac{x^4}{4} \pm \frac{x^5}{5} + \dots \quad \dots \quad (31)$$

and, if x is a small quantity, s , the exponential function, may be approximated by the equation,

$$e^{\pm s} = 1 \pm s; \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (32)$$

or, by the more general equation,

$$e^{\pm as} = 1 \pm as; \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (33)$$

and, if a greater accuracy is required, the second term may be included, thus,

$$e^{\pm s} = 1 \pm s + \frac{s^2}{2}, \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (34)$$

and then

$$e^{\pm as} = 1 \pm as + \frac{a^2 s^2}{2}. \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (35)$$

The logarithm is defined by $\log_e x = \int \frac{dx}{x}$; hence,

$$\log_e (1 \pm x) = \pm \int \frac{dx}{1 \pm x}.$$

Resolving $\frac{1}{1 \pm x}$ into a series, by (10), and then integrating, gives

$$\begin{aligned} \log_e (1 \pm x) &= \pm \int (1 \mp x + x^2 \mp x^3 + \dots) dx \\ &= \pm x - \frac{x^2}{2} \pm \frac{x^3}{3} - \frac{x^4}{4} \pm \frac{x^5}{5} - \dots \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (36) \end{aligned}$$

This logarithmic series (36) leads to the approximation,

$$\log_e (1 \pm s) = \pm s; \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (37)$$

or, including the second term, it gives

$$\log_e (1 \pm s) = \pm s - s^2, \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (38)$$

and the more general expression is, respectively,

$$\log_e (a \pm s) = \log a \left(1 \pm \frac{s}{a} \right) = \log a + \log \left(1 \pm \frac{s}{a} \right) = \log a \pm \frac{s}{a}, \quad (39)$$

and, more accurately,

$$\log_{\epsilon} (a \pm s) = \log a \pm \frac{s}{a} - \frac{s^2}{a^2} \cdot \cdot \cdot \cdot \quad (40)$$

Since $\log_{10} N = \log_{10} \epsilon \times \log_{\epsilon} N = 0.4343 \log_{\epsilon} N$, equations (39) and (40) may be written thus,

$$\left. \begin{aligned} \log_{10} (1 \pm s) &= \pm 0.4343s; \\ \log_{10} (a \pm s) &= \log_{10} a \pm 0.4343 \frac{s}{a} \end{aligned} \right\} \cdot \cdot \cdot \cdot \quad (41)$$

132. The trigonometric functions are represented by the infinite series:

$$\left. \begin{aligned} \cos x &= 1 - \frac{x^2}{2} + \frac{x^4}{4} - \frac{x^6}{6} + \dots; \\ \sin x &= x - \frac{x^3}{3} + \frac{x^5}{5} - \frac{x^7}{7} + \dots, \end{aligned} \right\} \cdot \cdot \cdot \cdot \quad (42)$$

which when s is a small quantity, may be approximated by

$$\cos s = 1 \quad \text{and} \quad \sin s = s; \quad \cdot \cdot \cdot \cdot \quad (43)$$

or, they may be represented in closer approximation by

$$\left. \begin{aligned} \cos s &= 1 - \frac{s^2}{2}; \\ \sin s &= s \left[-\frac{s^2}{6} \right]; \end{aligned} \right\} \cdot \cdot \cdot \cdot \quad (44)$$

or, by the more general expressions,

$$\left. \begin{aligned} \cos as &= 1 \quad \text{and} \quad \cos as = 1 - \frac{a^2 s^2}{2}; \\ \sin as &= as \quad \text{and} \quad \sin as = as \left[-\frac{a^2 s^2}{6} \right]. \end{aligned} \right\} \cdot \cdot \cdot \cdot \quad (45)$$

133. Other functions containing small terms may frequently be approximated by Taylor's series, or its special case, MacLaurin's series.

MacLaurin's series is written thus:

$$f(x) = f(0) + xf'(0) + \frac{x^2}{2} f''(0) + \frac{x^3}{3} f'''(0) + \dots, \cdot \quad (46)$$

where f' , f'' , f''' , etc., are respectively the first, second, third, etc., differential quotient of f ; hence,

$$\left. \begin{aligned} f(s) &= f(0) + sf'(0); \\ f(as) &= f(0) + asf'(0). \end{aligned} \right\} \dots \dots \dots (47)$$

Taylor's series is written thus,

$$f(b+x) = f(b) + xf'(b) + \frac{x^2}{2}f''(b) + \frac{x^3}{3}f'''(b) + \dots \dots \dots (48)$$

and leads to the approximations:

$$\left. \begin{aligned} f(b \pm s) &= f(b) \pm sf'(b); \\ f(b \pm as) &= f(b) \pm asf'(b). \end{aligned} \right\} \dots \dots \dots (49)$$

Many of the previously discussed approximations can be considered as special cases of (47) and (49).

134. As seen in the preceding, convenient equations for the approximation of expressions containing small terms are derived from various infinite series, which are summarized below:

$$\left. \begin{aligned} \frac{1}{1 \pm x} &= 1 \mp x + x^2 \mp x^3 + x^4 \mp \dots; \\ (1 \pm x)^n &= 1 \pm nx + \frac{n(n-1)}{2}x^2 \pm \frac{n(n-1)(n-2)}{3}x^3 + \dots; \\ e^{\pm x} &= 1 \pm x + \frac{x^2}{2} \pm \frac{x^3}{3} + \frac{x^4}{4} \pm \dots; \\ \log_e (1 \pm x) &= \pm x - \frac{x^2}{2} \pm \frac{x^3}{3} - \frac{x^4}{4} \pm \dots; \\ \cos x &= 1 - \frac{x^2}{2} + \frac{x^4}{4} - \frac{x^6}{6} + \dots; \\ \sin x &= x - \frac{x^3}{3} + \frac{x^5}{5} - \frac{x^7}{7} + \dots; \\ f(x) &= f(0) + xf'(0) + \frac{x^2}{2}f''(0) + \frac{x^3}{3}f'''(0) + \dots; \\ f(b \pm x) &= f(b) \pm xf'(b) + \frac{x^2}{2}f''(b) \pm \frac{x^3}{3}f'''(b) + \dots \end{aligned} \right\} (50)$$

The first approximations, derived by neglecting all higher terms but the first power of the small quantity $x=s$ in these series, are:

$$\left. \begin{aligned} \frac{1}{1 \pm s} &= 1 \mp s; & [+s^2]; \\ (1 \pm s)^n &= 1 \pm ns; & \left[+\frac{n(n-1)}{2}s^2 \right]; \\ e^{\pm s} &= 1 \pm s; & \left[+\frac{s^2}{2} \right]; \\ \log_e (1 \pm s) &= \pm s; & \left[-\frac{s^2}{2} \right]; \\ \cos s &= 1; & \left[-\frac{s^2}{2} \right]; \\ \sin s &= s; & \left[-\frac{s^3}{6} \right]; \\ f(s) &= f(0) + sf'(0); & \left[+\frac{s^2}{2}f''(0) \right]; \\ f(b \pm s) &= f(b) \pm sf'(b); & \left[+\frac{s^2}{2}f''(b) \right]; \end{aligned} \right\} \quad (51)$$

and, in addition hereto is to be remembered the multiplication rule,

$$(1 \pm s_1)(1 \pm s_2) = 1 \pm s_1 \pm s_2; \quad [\pm s_1 s_2]. \quad (52)$$

135. The accuracy of the approximation can be estimated by calculating the next term beyond that which is used. This term is given in brackets in the above equations (50) and (51).

Thus, when calculating a series of numerical values by approximation, for the one value, for which, as seen by the nature of the problem, the approximation is least close, the next term is calculated, and if this is less than the permissible limits of accuracy, the approximation is satisfactory.

For instance, in Example 2 of paragraph 130, the approximate value of the short-circuit current was found in (30), as

$$i = \frac{e_0}{x_0} \left\{ 1 - \frac{1}{2} \left(\frac{r}{ax_0} \right)^2 \right\}.$$

The next term in the parenthesis of equation (30), by the binomial, would have been $+\frac{n(n-1)}{2}s^2$; substituting $n=-\frac{1}{2}$; $s=\left(\frac{r}{ax_0}\right)^2$, the next becomes $+\frac{3}{8}\left(\frac{r}{ax_0}\right)^4$. The smaller the a , the less exact is the approximation.

The smallest value of a , considered in paragraph 130, was $a=0.2$. For $x_0=20r$, this gives $+\frac{3}{8}\left(\frac{r}{ax_0}\right)^4=0.00146$, as the value of the first neglected term, and in the accuracy of the result this is of the magnitude of $\frac{e_0}{x_0} \times 0.00146$, out of $\frac{e_0}{x_0} \times 0.9688$, the value given in paragraph 130; that is, the approximation gives the result correctly within $\frac{0.00146}{0.9688}=0.0015$ or within one-sixth of one per cent, which is sufficiently close for all engineering purposes, and with larger a the values are still closer approximations.

136. It is interesting to note the different expressions, which are approximated by $(1+s)$ and by $(1-s)$. Some of them are given in the following:

$1+s=$	$1-s=$
$\frac{1}{1-s};$	$\frac{1}{1+s};$
$\left(1+\frac{s}{n}\right)^n;$	$\left(1-\frac{s}{n}\right)^n;$
$\left(1+\frac{s}{2}\right)^2;$	$\left(1-\frac{s}{2}\right)^2;$
$\frac{1}{\left(1-\frac{s}{n}\right)^n};$	$\frac{1}{\left(1+\frac{s}{n}\right)^n};$
$\left(\frac{1+\frac{m}{n^2}s}{1-\frac{n-m}{n^2}s}\right)^n;$	$\left(\frac{1-\frac{m}{n^2}s}{1+\frac{n-m}{n^2}s}\right)^n;$

$$\frac{\sqrt{1+2s}}{1};$$

$$\frac{1}{\sqrt{1-2s}};$$

$$\sqrt{\frac{1+s}{1-s}};$$

$$\sqrt[n]{1+ns};$$

$$\frac{1}{\sqrt[n]{1-ns}};$$

$$\sqrt[n]{\frac{1+\frac{ns}{2}}{1-\frac{ns}{2}}};$$

$$\sqrt[n]{\frac{1+ms}{1-(n-m)s}};$$

etc.

$$\epsilon^s;$$

$$2-\epsilon^{-s};$$

$$1+\log_{\epsilon}(1+s);$$

$$1-\log_{\epsilon}(1-s);$$

$$1+n\log_{\epsilon}\left(1+\frac{s}{n}\right);$$

$$1-n\log_{\epsilon}\left(1-\frac{s}{n}\right);$$

$$1+\log_{\epsilon}\sqrt{\frac{1+s}{1-s}};$$

$$1-\log_{\epsilon}\sqrt{\frac{1-s}{1+s}};$$

etc.

$$1+\sin s;$$

$$1+n\sin\frac{s}{n};$$

$$\frac{\sqrt{1-2s}}{1};$$

$$\frac{1}{\sqrt{1+2s}};$$

$$\sqrt{\frac{1-s}{1+s}};$$

$$\sqrt[n]{1-ns};$$

$$\frac{1}{\sqrt[n]{1+ns}};$$

$$\sqrt[n]{\frac{1-\frac{ns}{2}}{1+\frac{ns}{2}}};$$

$$\sqrt[n]{\frac{1-ms}{1+(n-m)s}};$$

etc.

$$\epsilon^{-s};$$

$$2-\epsilon^s;$$

$$1+\log_{\epsilon}(1-s);$$

$$1-\log_{\epsilon}(1+s);$$

$$1+n\log_{\epsilon}\left(1-\frac{s}{n}\right);$$

$$1-n\log_{\epsilon}\left(1+\frac{s}{n}\right);$$

$$1+\log_{\epsilon}\sqrt{\frac{1-s}{1+s}};$$

$$1-\log_{\epsilon}\sqrt{\frac{1+s}{1-s}};$$

etc.

$$1-\sin s;$$

$$1-n\sin\frac{s}{n};$$

$$\begin{array}{ll}
 1 + \frac{1}{n} \sin ns; & 1 - \frac{1}{n} \sin ns; \\
 \cos \sqrt{-2s}; & \cos \sqrt{2s}; \\
 \text{etc.} & \text{etc.}
 \end{array}$$

137. As an example may be considered the reduction to its simplest form, of the expression:

$$F = \frac{\sqrt{a} \sqrt{(a+s_1)^3} \{4 - \sin 6s_2\} \sqrt[4]{a} \epsilon^{\frac{2s_1}{a}} \cos^2 \sqrt{\frac{2s_1}{a}}}{\epsilon^{-3s_2}(a+2s_1) \left\{1 - a \log \epsilon \sqrt{\frac{a-s_2}{a+s_2}}\right\} \sqrt{a-2s_1}};$$

then,

$$\sqrt{(a+s_1)^3} = (a+s_1)^{3/4} = a^{3/4} \left(1 + \frac{s_1}{a}\right)^{3/4} = a^{3/4} \left(1 + \frac{3}{4} \frac{s_1}{a}\right);$$

$$4 - \sin 6s_2 = 4 \left(1 - \frac{1}{4} \sin 6s_2\right) = 4 \left(1 - \frac{3}{2} s_2\right);$$

$$\epsilon^{\frac{2s_1}{a}} = 1 + 2 \frac{s_1}{a};$$

$$\cos^2 \sqrt{\frac{2s_1}{a}} = \left(1 - \frac{s_1}{a}\right)^2 = 1 - 2 \frac{s_1}{a};$$

$$\epsilon^{-3s_2} = 1 - 3s_2;$$

$$a + 2s_1 = a \left(1 + 2 \frac{s_1}{a}\right);$$

$$\begin{aligned}
 1 - a \log \epsilon \sqrt{\frac{a-s_2}{a+s_2}} &= 1 - a \log \epsilon \sqrt{\frac{1 - \frac{s_2}{a}}{1 + \frac{s_2}{a}}} = 1 - a \log \epsilon \sqrt{1 - 2 \frac{s_2}{a}} \\
 &= 1 - a \log \epsilon \left(1 - \frac{s_2}{a}\right) = 1 + s_2;
 \end{aligned}$$

$$\sqrt{a-2s_1} = a^{1/2} \left(1 - \frac{2s_1}{a}\right)^{1/2} = a^{1/2} \left(1 - \frac{s_1}{a}\right);$$

hence,

$$\begin{aligned}
 F &= \frac{a^{1/2} \times a^{3/4} \left(1 + \frac{3}{4} \frac{s_1}{a}\right) \times 4 \left(1 - \frac{3}{2} s_2\right) \times a^{1/4} \times \left(1 + 2 \frac{s_1}{a}\right) \left(1 - 2 \frac{s_1}{a}\right)}{(1 - 3s_2) \times a \left(1 + 2 \frac{s_1}{a}\right) (1 + s_2) \times a^{1/2} \left(1 - \frac{s_1}{a}\right)} \\
 &= \frac{4a^{3/2} \left(1 + \frac{3}{4} \frac{s_1}{a} - \frac{3}{2} s_2 + 2 \frac{s_1}{a} - 2 \frac{s_1}{a}\right)}{a^{3/2} \left(1 - 3s_2 + 2 \frac{s_1}{a} + s_2 - \frac{s_1}{a}\right)} \\
 &= 4 \left(1 - \frac{1}{4} \frac{s_1}{a} + \frac{s_2}{2a}\right).
 \end{aligned}$$

138. As further example may be considered the equations of an alternating-current electric circuit, containing distributed resistance, inductance, capacity, and shunted conductance, for instance, a long-distance transmission line or an underground high-potential cable.

Equations of the Transmission Line.

Let l be the distance along the line, from some starting point; E , the voltage; I , the current at point l , expressed as vector quantities or general numbers; $Z_0 = r_0 + jx_0$, the line impedance per unit length (for instance, per mile); $Y_0 = g_0 + jb_0$ = line admittance, shunted, per unit length; that is, r_0 is the ohmic effective resistance; x_0 , the self-inductive reactance; b_0 , the condensive susceptance, that is, wattless charging current divided by volts, and g_0 = energy component of admittance, that is, energy component of charging current, divided by volts, per unit length, as, per mile.

Considering a line element dl , the voltage, dE , consumed by the impedance is $Z_0 I dl$, and the current, dI , consumed by the admittance is $Y_0 E dl$; hence, the following relations may be written:

$$\frac{dE}{dl} = Z_0 I; \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (1)$$

$$\frac{dI}{dl} = Y_0 E. \quad . \quad . \quad . \quad . \quad . \quad . \quad (2)$$

Differentiating (1), and substituting (2) therein gives

$$\frac{d^2 \dot{E}}{dl^2} = Z_0 Y_0 \dot{E}, \quad . \quad . \quad . \quad . \quad . \quad (3)$$

and from (1) it follows that,

$$I = \frac{1}{Z_0} \frac{d\dot{E}}{dl}. \quad . \quad . \quad . \quad . \quad . \quad (4)$$

Equation (3) is integrated by

$$\dot{E} = A e^{Bl}, \quad . \quad . \quad . \quad . \quad . \quad (5)$$

and (5) substituted in (3) gives

$$B = \pm \sqrt{Z_0 Y_0}; \quad . \quad . \quad . \quad . \quad . \quad (6)$$

hence, from (5) and (4), it follows

$$\dot{E} = A_1 e^{+\sqrt{Z_0 Y_0} l} + A_2 e^{-\sqrt{Z_0 Y_0} l}; \quad . \quad . \quad . \quad . \quad . \quad (7)$$

$$I = \sqrt{\frac{Y_0}{Z_0}} \{A_1 e^{+\sqrt{Z_0 Y_0} l} - A_2 e^{-\sqrt{Z_0 Y_0} l}\}. \quad . \quad . \quad . \quad . \quad (8)$$

Next assume

$$\left. \begin{aligned} l &= l_0, \quad \text{the entire length of line;} \\ Z &= l_0 Z_0, \text{ the total line impedance;} \\ \text{and } Y &= l_0 Y_0, \text{ the total line admittance;} \end{aligned} \right\} \quad . \quad . \quad . \quad (9)$$

then, substituting (9) into (7) and (8), the following expressions are obtained:

$$\left. \begin{aligned} \dot{E}_1 &= A_1 e^{+\sqrt{ZY}} + A_2 e^{-\sqrt{ZY}}; \\ I_1 &= \sqrt{\frac{Y}{Z}} \{A_1 e^{+\sqrt{ZY}} - A_2 e^{-\sqrt{ZY}}\}, \end{aligned} \right\} \quad . \quad . \quad . \quad (10)$$

as the voltage and current at the generator end of the line.

139. If now $\dot{E}_0$ and I_0 respectively are the current and voltage at the step-down end of the line, for $l=0$, by substituting $l=0$ into (7) and (8),

$$\left. \begin{aligned} A_1 + A_2 &= \dot{E}_0; \\ A_1 - A_2 &= I_0 \sqrt{\frac{Z}{Y}}. \end{aligned} \right\} \quad . \quad . \quad . \quad . \quad (11)$$

Substituting in (10) for the exponential function, the series,

$$\begin{aligned} e^{\pm ZY} &= 1 \pm \sqrt{ZY} + \frac{ZY}{2} \pm \frac{ZY\sqrt{ZY}}{6} + \frac{Z^2Y^2}{24} \pm \frac{Z^2Y^2\sqrt{ZY}}{120} + \dots \\ &= \left(1 + \frac{ZY}{2} + \frac{Z^2Y^2}{24}\right) \pm \sqrt{ZY} \left(1 + \frac{ZY}{6} + \frac{Z^2Y^2}{120}\right), \quad \dots \quad (12) \end{aligned}$$

and arranging by $(A_1 + A_2)$ and $(A_1 - A_2)$, and substituting herefor the expressions (11), gives

$$\begin{aligned} E_1 &= E_0 \left\{ 1 + \frac{ZY}{2} + \frac{Z^2Y^2}{24} \right\} + ZI_0 \left\{ 1 + \frac{ZY}{6} + \frac{Z^2Y^2}{120} \right\}; \\ I_1 &= I_0 \left\{ 1 + \frac{ZY}{2} + \frac{Z^2Y^2}{24} \right\} + YE_0 \left\{ 1 + \frac{ZY}{6} + \frac{Z^2Y^2}{120} \right\}. \end{aligned} \quad (13)$$

When $l = -l_0$, that is, for E_0 and I_0 at the generator side, and E_1 and I_1 at the step-down side of the line, the sign of the second term of equations (13) merely reverses.

140. From the foregoing, it follows that, if Z is the total impedance; Y , the total shunted admittance of a transmission line; E_0 and I_0 , the voltage and current at one end; E_1 and I_1 , the voltage and current at the other end of the transmission line; then,

$$\begin{aligned} E_1 &= E_0 \left\{ 1 + \frac{ZY}{2} + \frac{Z^2Y^2}{24} \right\} \pm ZI_0 \left\{ 1 + \frac{ZY}{6} + \frac{Z^2Y^2}{120} \right\}; \\ I_1 &= I_0 \left\{ 1 + \frac{ZY}{2} + \frac{Z^2Y^2}{24} \right\} \pm YE_0 \left\{ 1 + \frac{ZY}{6} + \frac{Z^2Y^2}{120} \right\}; \end{aligned} \quad (14)$$

where the plus sign applies if E_0, I_0 is the step-down end, the minus sign, if E_0, I_0 is the step-up end of the transmission line.

In practically all cases, the quadratic term can be neglected, and the equations simplified, thus,

$$\begin{aligned} E_1 &= E_0 \left\{ 1 + \frac{ZY}{2} \right\} \pm ZI_0 \left\{ 1 + \frac{ZY}{6} \right\}; \\ I_1 &= I_0 \left\{ 1 + \frac{ZY}{2} \right\} \pm YE_0 \left\{ 1 + \frac{ZY}{6} \right\}, \end{aligned} \quad (15)$$

and the error made hereby is of the magnitude of less than $\frac{z^2y^2}{24}$.

Except in the case of very long lines, the second term of the second term can also usually be neglected, which gives

$$\left. \begin{aligned} E_1 &= E_0 \left(1 + \frac{ZY}{2} \right) \pm Z I_0; \\ I_1 &= I_0 \left(1 + \frac{ZY}{2} \right) \pm Y E_0. \end{aligned} \right\} \dots \dots (16)$$

and the error made hereby is of the magnitude of less than $\frac{zy}{6}$ of the line impedance voltage and line charging current.

141. Example. Assume 200 miles of 60-cycle line, on non-inductive load of $e_0 = 100,000$ volts; and $i_0 = 100$ amperes. The line constants, as taken from tables are $Z = 104 + 140j$ ohms and $Y = +0.0013j$ ohms; hence,

$$ZY = -(0.182 - 0.136j);$$

$$\begin{aligned} E_1 &= 100000(1 - 0.091 + 0.068j) + 100(104 + 140j) \\ &= 101400 + 20800j, \text{ in volts;} \end{aligned}$$

$$\begin{aligned} I_1 &= 100(1 - 0.091 + 0.068j) + 0.0013j \times 100000 \\ &= 91 + 136.8j, \text{ in amperes.} \end{aligned}$$

$$\text{The error is } \frac{zy}{6} = \frac{0.174 \times 0.0013}{6} = \frac{0.226}{6} = 0.038.$$

In E_1 , the neglect of the second term of $zI_0 = 17,400$, gives an error of $0.038 \times 17,400 = 660$ volts = 0.6 per cent.

In I_1 , the neglect of the second term of $yE_0 = 130$, gives an error of $0.038 \times 130 = 5$ amperes = 3 per cent.

Although the charging current of the line is 130 per cent of output current, the error in the current is only 3 per cent.

Using the equations (15), which are nearly as simple, brings the error down to $\frac{z^2 y^2}{24} = \frac{0.226^2}{24} = 0.0021$, or less than one-quarter per cent.

Hence, only in extreme cases the equations (14) need to be used. Their error would be less than $\frac{z^4 y^4}{720} = 3.6 \times 10^{-6}$, or one three-thousandth per cent.

The accuracy of the preceding approximation can be estimated by considering the physical meaning of Z and Y : Z is the line impedance; hence ZI the impedance voltage, and $u = \frac{ZI}{E}$, the impedance voltage of the line, as fraction of total voltage; Y is the shunted admittance; hence YE the charging current, and $v = \frac{YE}{I}$, the charging current of the line, as fraction of total current.

Multiplying gives $uv = ZY$; that is, the constant ZY is the product of impedance voltage and charging current, expressed as fractions of full voltage and full current, respectively. In any economically feasible power transmission, irrespective of its length, both of these fractions, and especially the first, must be relatively small, and their product therefore is a small quantity, and its higher powers negligible.

In any economically feasible constant potential transmission line the preceding approximations are therefore permissible.

Approximation by Chain Fraction.

141A.—A convenient method of approximating numerical values is often afforded by the chain fraction. A chain fraction is a fraction, in which the denominator contains a fraction, which again in its denominator contains a fraction, etc. Thus:

$$\frac{1}{2 + \frac{1}{3 + \frac{1}{1 + \frac{1}{4}}}}$$

Only integer chain fractions, that is, chain fractions in which all numerators are unity, are of interest.

A common fraction is converted into a chain fraction thusly:

$$\frac{511}{1152} = ?$$

$$\frac{511}{1152} = \frac{1}{\frac{1152}{511}} = \frac{1}{2 + \frac{130}{511}}$$

$$= \frac{1}{2 + \frac{1}{\frac{511}{130}}} = \frac{1}{2 + \frac{1}{3 + \frac{121}{130}}}$$

$$= \frac{1}{2 + \frac{1}{3 + \frac{1}{\frac{130}{121}}}} = \frac{1}{2 + \frac{1}{3 + \frac{1}{1 + \frac{9}{121}}}}$$

$$= \frac{1}{2 + \frac{1}{3 + \frac{1}{1 + \frac{1}{\frac{121}{9}}}}}} = \frac{1}{2 + \frac{1}{3 + \frac{1}{1 + \frac{1}{13 + \frac{4}{9}}}}}}$$

$$= \frac{1}{2 + \frac{1}{3 + \frac{1}{1 + \frac{1}{13 + \frac{1}{9}}}}}} = \frac{1}{2 + \frac{1}{3 + \frac{1}{1 + \frac{1}{13 + \frac{1}{2 + \frac{1}{4}}}}}}}}$$

That is, to convert a common fraction into a chain fraction, the numerator is divided into the denominator, the residue divided into the divisor, and so on, until no residue remains. The successive quotients then are the successive denominators of the chain fraction.

For instance:

$$\frac{511}{1152} = ?$$

$$511/1152 = 2$$

$$\frac{1022}{130/511 = 3}$$

$$\frac{390}{121/130 = 1}$$

$$\frac{121}{9/121 = 13}$$

$$\frac{9}{9}$$

$$\frac{31}{27}$$

$$\frac{4/9 = 2}{8}$$

$$\frac{1/4 = 4}{1/4 = 4}$$

$$\frac{4/9 = 2}{8}$$

$$\frac{8}{1/4 = 4}$$

$$\frac{1/4 = 4}{1/4 = 4}$$

$$\frac{1/4 = 4}{1/4 = 4}$$

hence:

$$\frac{511}{1152} = \frac{1}{2 + \frac{1}{3 + \frac{1}{1 + \frac{1}{13 + \frac{1}{2 + \frac{1}{\frac{1}{4}}}}}}}$$

Inversely, the chain fraction is converted into a common fraction, by rolling it up from the end:

$$2 + \frac{1}{4} = \frac{9}{4}$$

$$\frac{1}{2 + \frac{1}{4}} = \frac{4}{9}$$

$$13 + \frac{1}{2 + \frac{1}{4}} = \frac{121}{9}$$

$$\frac{1}{13 + \frac{1}{2 + \frac{1}{4}}} = \frac{9}{121}$$

$$1 + \frac{1}{13 + \frac{1}{2 + \frac{1}{4}}} = \frac{130}{121}$$

$$\frac{1}{1 + \frac{1}{13 + \frac{1}{2 + \frac{1}{4}}}} = \frac{121}{130}$$

$$3 + \frac{1}{1 + \frac{1}{13 + \frac{1}{2 + \frac{1}{4}}}} = \frac{511}{130}$$

$$\frac{1}{3 + \frac{1}{1 + \frac{1}{13 + \frac{1}{2 + \frac{1}{4}}}}} = \frac{130}{511}$$

$$2 + \frac{1}{3 + \frac{1}{1 + \frac{1}{13 + \frac{1}{2 + \frac{1}{4}}}}} = \frac{1152}{511}$$

$$\frac{1}{2 + \frac{1}{3 + \frac{1}{1 + \frac{1}{13 + \frac{1}{2 + \frac{1}{4}}}}} = \frac{511}{1152}$$

The expression of the numerical value by chain fraction gives a series of successive approximations. Thus the successive approximation of the chain fraction:

$$\frac{1}{2 + \frac{1}{3 + \frac{1}{1 + \frac{1}{13 + \frac{1}{2 + \frac{1}{4}}}}}} = \frac{511}{1152}$$

are:

difference: = %:

$$(1) \frac{1}{2} = \frac{1}{2} = .5 \quad \dots + .0564 = +12.7\%$$

$$(2) \frac{1}{2 + \frac{1}{3}} = \frac{3}{7} = .42857 \quad \dots - .0150 = -3.4\%$$

$$(3) \frac{1}{2 + \frac{1}{3 + \frac{1}{1}}} = \frac{4}{9} = .44444 \quad \dots + .00086 = +.194\%$$

$$(4) \frac{1}{2 + \frac{1}{3 + \frac{1}{1 + \frac{1}{13}}}} = \frac{55}{124} = .443548 \quad \dots - .000028 = -.0068\%$$

$$(5) \frac{1}{2 + \frac{1}{3 + \frac{1}{1 + \frac{1}{13 + \frac{1}{2}}}}} = \frac{114}{257} = .443580 \dots + .000004 = +.0009\%$$

$$(6) \frac{1}{2 + \frac{1}{3 + \frac{1}{1 + \frac{1}{13 + \frac{1}{2 + \frac{1}{4}}}}}} = \frac{511}{1152} = .443576 \dots$$

As seen, successive approximations are alternately above and below the true value, and the approach to the true value is extremely rapid. It is the latter feature which makes the chain fraction valuable, as where it can be used, it gives very rapidly converging approximations.

141B.—Chain fraction representing irrational numbers, as π , e , etc., may be endless. Thus:

$$\pi = 3 + \frac{1}{7 + \frac{1}{15 + \frac{1}{1 + \frac{1}{288 + \frac{1}{1 + \frac{1}{2 + \frac{1}{1 + \frac{1}{3 + \frac{1}{1 + \frac{1}{7 + \dots}}}}}}}}}} = 3.14159265 \dots$$

The first three approximations of this chain fraction of π are:

difference: = %

$$(1) 3 + \frac{1}{7} = 3 \frac{1}{7} = 3.142857 \dots + .00127 = + .043\%$$

$$(2) 3 + \frac{1}{7 + \frac{1}{15}} = 3 \frac{15}{106} = 3.1415094 \dots - .0000832 = - .0026\%$$

$$(3) 3 + \frac{1}{7 + \frac{1}{15 + \frac{1}{1}}} = 3 \frac{16}{113} = 3.1415929 \dots + .0000003 = + .000009\%$$

As seen, the first approximation, $3 \frac{1}{7}$, is already sufficiently close for most practical purposes, and the third approximation of the chain fraction is correct to the 6th decimal.

144.—Frequently irrational numbers, such as square roots, can be expressed by periodic chain fractions, and the chain

fraction offers a convenient way of expressing numerical values containing square roots, and deriving their approximations.

For instance:

Resolve $\sqrt{6}$ into a chain fraction.

As the chain fraction is <1 , $\sqrt{6}$ has to be expressed in the form:

$$\sqrt{6} = 2 + (\sqrt{6} - 2) \quad (1)$$

and the latter term: $(\sqrt{6} - 2)$, which is <1 , expressed as chain fraction.

To rationalize the numerator, we multiply numerator and denominator by $(\sqrt{6} + 2)$:

$$(\sqrt{6} - 2) = \frac{(\sqrt{6} - 2)(\sqrt{6} + 2)}{\sqrt{6} + 2} = \frac{2}{\sqrt{6} + 2} = \frac{1}{\frac{\sqrt{6} + 2}{2}}$$

thus:

$$\sqrt{6} = 2 + \frac{1}{\frac{\sqrt{6} + 2}{2}}$$

as $\frac{\sqrt{6} + 2}{2}$ is > 1 , it is again resolved into:

$$\frac{\sqrt{6} + 2}{2} = 2 + \frac{\sqrt{6} - 2}{2}$$

thus:

$$\sqrt{6} = 2 + \frac{1}{2 + \frac{\sqrt{6} - 2}{2}}$$

continuing in the same manner:

$$\frac{\sqrt{6} - 2}{2} = \frac{(\sqrt{6} - 2)(\sqrt{6} + 2)}{2(\sqrt{6} + 2)} = \frac{2}{2(\sqrt{6} + 2)} = \frac{1}{\sqrt{6} + 2}$$

hence:

$$\sqrt{6} = 2 + \frac{1}{2 + \frac{1}{\sqrt{6} + 2}}$$

and:

$$\sqrt{6} + 2 = 4 + (\sqrt{6} - 2)$$

hence:

$$\sqrt{6} = 2 + \frac{1}{2 + \frac{1}{4 + (\sqrt{6} - 2)}}$$

and, as the term $(\sqrt{6} - 2)$ appeared already at (1), we are here at the end of the recurring period, that is, the denominators now repeat:

$$\sqrt{6} = 2 + \frac{1}{2 + \frac{1}{4 + \frac{1}{2 + \frac{1}{4 + \frac{1}{2 + \dots}}}}}$$

a periodic chain fraction, in which the denominators 2 and 4 alternate.

In the same manner,

$$\sqrt{2} = 1 + \frac{1}{2 + \frac{1}{2 + \frac{1}{2 + \dots}}} \quad \text{with the periodic denominator 2}$$

$$\sqrt{3} = 1 + \frac{1}{1 + \frac{1}{2 + \frac{1}{1 + \frac{1}{2 + \dots}}}} \quad \text{with the periodic denominators 1 and 2}$$

$$\sqrt{5} = 2 + \frac{1}{4 + \frac{1}{4 + \frac{1}{4 + \dots}}} \quad \text{with the periodic denominator 4}$$

This method of resolution of roots into chain fractions gives a convenient way of deriving simple numerical approximations of the roots, and hereby is very useful.

For instance, the third approximation of $\sqrt{2}$ is $1\frac{5}{12}$, with an error of .2 per cent, that is, close enough for most practical purposes. Thus, the diagonal of a square with 1 foot as side, is very closely 1 foot 5 inches, etc.



CHAPTER VI.

EMPIRICAL CURVES.

A. General.

142. The results of observation or tests usually are plotted in a curve. Such curves, for instance, are given by the core loss of an electric generator, as function of the voltage; or, the current in a circuit, as function of the time, etc. When plotting from numerical observations, the curves are empirical, and the first and most important problem which has to be solved to make such curves useful is to find equations for the same, that is, find a function, $y=f(x)$, which represents the curve. As long as the equation of the curve is not known its utility is very limited. While numerical values can be taken from the plotted curve, no general conclusions can be derived from it, no general investigations based on it regarding the conditions of efficiency, output, etc. An illustration hereof is afforded by the comparison of the electric and the magnetic circuit. In the electric circuit, the relation between e.m.f. and

current is given by Ohm's law, $i = \frac{e}{r}$, and calculations are uni-

versally and easily made. In the magnetic circuit, however, the term corresponding to the resistance, the reluctance, is not a constant, and the relation between m.m.f. and magnetic flux cannot be expressed by a general law, but only by an empirical curve, the magnetic characteristic, and as the result, calculations of magnetic circuits cannot be made as conveniently and as general in nature as calculations of electric circuits.

If by observation or test a number of corresponding values of the independent variable x and the dependent variable y are determined, the problem is to find an equation, $y=f(x)$, which represents these corresponding values: $x_1, x_2, x_3 \dots x_n$, and $y_1, y_2, y_3 \dots y_n$, approximately, that is, within the errors of observation.

The mathematical expression which represents an empirical curve may be a rational equation or an empirical equation. It is a rational equation if it can be derived theoretically as a conclusion from some general law of nature, or as an approximation thereof, but it is an empirical equation if no theoretical reason can be seen for the particular form of the equation. For instance, when representing the dying out of an electrical current in an inductive circuit by an exponential function of time, we have a rational equation: the induced voltage, and therefore, by Ohm's law, the current, varies proportionally to the rate of change of the current, that is, its differential quotient, and as the exponential function has the characteristic of being proportional to its differential quotient, the exponential function thus rationally represents the dying out of the current in an inductive circuit. On the other hand, the relation between the loss by magnetic hysteresis and the magnetic density: $W = \eta B^{1.6}$, is an empirical equation since no reason can be seen for this law of the 1.6th power, except that it agrees with the observations.

A rational equation, as a deduction from a general law of nature, applies universally, within the range of the observations as well as beyond it, while an empirical equation can with certainty be relied upon only within the range of observation from which it is derived, and extrapolation beyond this range becomes increasingly uncertain. A rational equation therefore is far preferable to an empirical one. As regards the accuracy of representing the observations, no material difference exists between a rational and an empirical equation. An empirical equation frequently represents the observations with great accuracy, while inversely a rational equation usually does not rigidly represent the observations, for the reason that in nature the conditions on which the rational law is based are rarely perfectly fulfilled. For instance, the representation of a decaying current by an exponential function is based on the assumption that the resistance and the inductance of the circuit are constant, and capacity absent, and none of these conditions can ever be perfectly satisfied, and thus a deviation occurs from the theoretical condition, by what is called "secondary effects."

143. To derive an equation, which represents an empirical curve, careful consideration should first be given to the physical

nature of the phenomenon which is to be expressed, since thereby the number of expressions which may be tried on the empirical curve is often greatly reduced. Much assistance is usually given by considering the zero points of the curve and the points at infinity. For instance, if the observations represent the core loss of a transformer or electric generator, the curve must go through the origin, that is, $y=0$ for $x=0$, and the mathematical expression of the curve $y=f(x)$ can contain no constant term. Furthermore, in this case, with increasing x , y must continuously increase, so that for $x=\infty$, $y=\infty$. Again, if the observations represent the dying out of a current as function of the time, it is obvious that for $x=\infty$, $y=0$. In representing the power consumed by a motor when running without load, as function of the voltage, for $x=0$, y cannot be $=0$, but must equal the mechanical friction, and an expression like $y=Ax^2$ cannot represent the observations, but the equation must contain a constant term.

Thus, first, from the nature of the phenomenon, which is represented by the empirical curve, it is determined

(a) Whether the curve is periodic or non-periodic.

(b) Whether the equation contains constant terms, that is, for $x=0$, $y \neq 0$, and inversely, or whether the curve passes through the origin: that is, $y=0$ for $x=0$, or whether it is hyperbolic; that is, $y = \infty$ for $x=0$, or $x=\infty$ for $y=0$.

(c) What values the expression reaches for ∞ . That is, whether for $x=\infty$, $y=\infty$, or $y=0$, and inversely.

(d) Whether the curve continuously increases or decreases, or reaches maxima and minima.

(e) Whether the law of the curve may change within the range of the observations, by some phenomenon appearing in some observations which does not occur in the other. Thus, for instance, in observations in which the magnetic density enters, as core loss, excitation curve, etc., frequently the curve law changes with the beginning of magnetic saturation, and in this case only the data below magnetic saturation would be used for deriving the theoretical equations, and the effect of magnetic saturation treated as secondary phenomenon. Or, for instance, when studying the excitation current of an induction motor, that is, the current consumed when running light, at low voltage the current may increase again with decreasing voltage.

instead of decreasing, as result of the friction load, when the voltage is so low that the mechanical friction constitutes an appreciable part of the motor output. Thus, empirical curves can be represented by a single equation only when the physical conditions remain constant within the range of the observations.

From the shape of the curve then frequently, with some experience, a guess can be made on the probable form of the equation which may express it. In this connection, therefore, it is of the greatest assistance to be familiar with the shapes of the more common forms of curves, by plotting and studying various forms of equations $y=f(x)$.

By changing the scale in which observations are plotted the apparent shape of the curve may be modified, and it is therefore desirable in plotting to use such a scale that the average slope of the curve is about 45 deg. A much greater or much lesser slope should be avoided, since it does not show the character of the curve as well.

B. Non-Periodic Curves.

144. The most common non-periodic curves are the potential series, the parabolic and hyperbolic curves, and the exponential and logarithmic curves.

THE POTENTIAL SERIES.

Theoretically, any set of observations can be represented exactly by a potential series of any one of the following forms:

$$y = a_0 + a_1x + a_2x^2 + a_3x^3 + \dots; \quad . \quad . \quad . \quad . \quad (1)$$

$$y = a_1x + a_2x^2 + a_3x^3 + \dots; \quad . \quad . \quad . \quad . \quad (2)$$

$$y = a_0 + \frac{a_1}{x} + \frac{a_2}{x^2} + \frac{a_3}{x^3} + \dots; \quad . \quad . \quad . \quad . \quad (3)$$

$$y = \frac{a_1}{x} + \frac{a_2}{x^2} + \frac{a_3}{x^3} + \dots \quad . \quad . \quad . \quad . \quad (4)$$

if a sufficiently large number of terms are chosen.

For instance, if n corresponding numerical values of x and y are given, $x_1, y_1; x_2, y_2; \dots x_n, y_n$, they can be represented

by the series (1), when choosing as many terms as required to give n constants a :

$$y = a_0 + a_1x + a_2x^2 + \dots + a_{n-1}n^{n-1}. \quad (5)$$

By substituting the corresponding values $x_1, y_1; x_2, y_2, \dots$ into equation (5), there are obtained n equations, which determine the n constants $a_0, a_1, a_2, \dots a_{n-1}$.

Usually, however, such representation is irrational, and therefore meaningless and useless.

TABLE I.

$\frac{e}{100} = x$	$P_i = y$	-0.5	+2x	+2.5x ²	-1.5x ³	+1.5x ⁴	-2x ⁵	+x ⁶
0.4	0.63	-0.5	+0.8	+0.4	-0.10	+0.04	-0.02	0
0.6	1.36	-0.5	+1.2	+0.9	-0.32	+0.19	-0.16	+0.05
0.8	2.18	-0.5	+1.6	+1.6	-0.77	+0.61	-0.65	+0.26
1.0	3.00	-0.5	+2.0	+2.5	-1.50	+1.50	-2.00	+1.00
1.2	3.93	-0.5	+2.4	+3.6	-2.59	+3.11	-4.98	+2.89
1.4	6.22	-0.5	+2.8	+4.9	-4.12	+5.76	-10.76	+6.13
1.6	8.59	-0.5	+3.2	+6.4	-6.14	+9.83	-20.97	+16.78

Let, for instance, the first column of Table I represent the voltage, $\frac{e}{100} = x$, in hundreds of volts, and the second column the core loss, $p_i = y$, in kilowatts, of an 125-volt 100-h.p. direct-current motor. Since seven sets of observations are given, they can be represented by a potential series with seven constants, thus,

$$y = a_0 + a_1x + a_2x^2 + \dots + a_6x^6, \quad (6)$$

and by substituting the observations in (6), and calculating the constants a from the seven equations derived in this manner, there is obtained as empirical expression of the core loss of the motor the equation,

$$y = -0.5 + 2x + 2.5x^2 - 1.5x^3 + 1.5x^4 - 2x^5 + x^6. \quad (7)$$

This expression (7), however, while exactly representing the seven observations, has no physical meaning, as easily seen by plotting the individual terms. In Fig. 60, y appears

as the resultant of a number of large positive and negative terms. Furthermore, if one of the observations is omitted, and the potential series calculated from the remaining six values, a series reaching up to x^5 would be the result, thus,

$$y = a_0 + a_1x + a_2x^2 + \dots + a_5x^5, \quad \dots \quad (8)$$

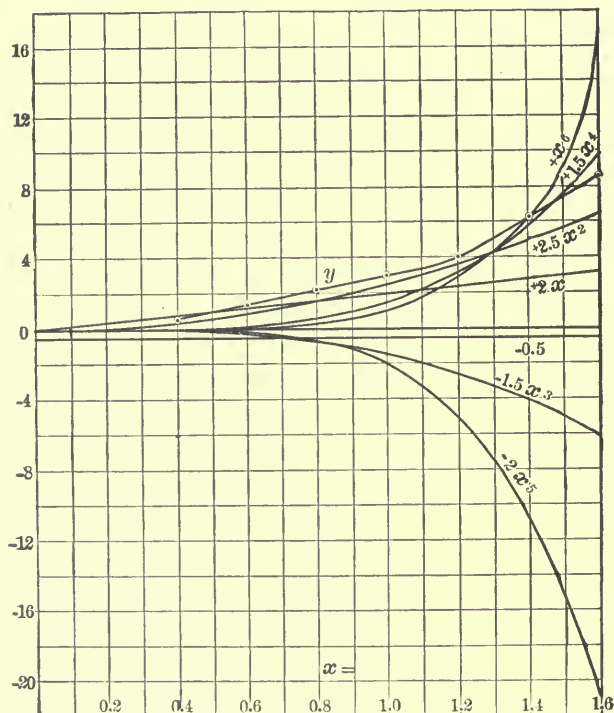


FIG. 60. Terms of Empirical Expression of Excitation Power.

but the constants a in (8) would have entirely different numerical values from those in (7), thus showing that the equation (7) has no rational meaning.

145. The potential series (1) to (4) thus can be used to represent an empirical curve only under the following conditions:

1. If the successive coefficients $a_0, a_1, a_2, \dots$ decrease in value so rapidly that within the range of observation the higher terms become rapidly smaller and appear as mere secondary terms.

2. If the successive coefficients a follow a definite law, indicating a convergent series which represents some other function, as an exponential, trigonometric, etc.

3. If all the coefficients, a , are very small, with the exception of a few of them, and only the latter ones thus need to be considered.

TABLE II.

x	y	y'	y_1
0.4	0.89	0.88	0.01
0.6	1.35	1.34	0.01
0.8	1.96	1.94	0.02
1.0	2.72	2.70	0.02
1.2	3.62	2.59	0.03
1.4	4.63	4.59	0.04
1.6	5.76	5.65	0.11

For instance, let the numbers in column 1 of Table II represent the speed x of a fan motor, as fraction of the rated speed, and those in column 2 represent the torque y , that is, the turning moment of the motor. These values can be represented by the equation,

$$y = 0.5 + 0.02x + 2.5x^2 - 0.3x^3 + 0.015x^4 - 0.02x^5 + 0.01x^6. \quad (9)$$

In this case, only the constant term and the terms with x^2 and x^3 have appreciable values, and the remaining terms probably are merely the result of errors of observations, that is, the approximate equation is of the form,

$$y = a_0 + a_2x^2 + a_3x^3. \quad (10)$$

Using the values of the coefficients from (9), gives

$$y = 0.5 + 2.5x^2 - 0.3x^3. \quad (11)$$

The numerical values calculated from (11) are given in column 3 of Table II as y' , and the difference between them and the observations of column 2 are given in column 4, as y_1 .

The values of column 4 can now be represented by the same form of equation, namely,

$$y_1 = b_0 + b_2x^2 + b_3x^3, \quad . \quad . \quad . \quad . \quad . \quad . \quad (12)$$

in which the constants b_0 , b_2 , b_3 are calculated by the method of least squares, as described in paragraph 120 of Chapter IV, and give

$$y_1 = 0.031 - 0.093x^2 + 0.076x^3. \quad . \quad . \quad . \quad . \quad (13)$$

Equation (13) added to (11) gives the final approximate equation of the torque, as,

$$y_0 = 0.531 + 2.407x^2 - 0.224x^3. \quad . \quad . \quad . \quad . \quad (14)$$

The equation (14) probably is the approximation of a rational equation, since the first term, 0.531, represents the bearing friction; the second term, $2.407x^2$ (which is the largest), represents the work done by the fan in moving the air, a resistance proportional to the square of the speed, and the third term approximates the decrease of the air resistance due to the churning motion of the air created by the fan.

In general, the potential series is of limited usefulness; it rarely has a rational meaning and is mainly used, where the curve approximately follows a simple law, as a straight line, to represent by small terms the deviation from this simple law, that is, the secondary effects, etc. Its use, thus, is often temporary, giving an empirical approximation pending the derivation of a more rational law.

The Parabolic and the Hyperbolic Curves.

146. One of the most useful classes of curves in engineering are those represented by the equation,

$$y = ax^n; \quad . \quad . \quad . \quad . \quad . \quad . \quad (15)$$

or, the more general equation,

$$y - b = a(x - c)^n. \quad . \quad . \quad . \quad . \quad . \quad (16)$$

Equation (16) differs from (15) only by the constant terms b and c ; that is, it gives a different location to the coordinate

center, but the curve shape is the same, so that in discussing the general shapes, only equation (15) need be considered.

If n is positive, the curves $y=ax^n$ are *parabolic curves*, passing through the origin and increasing with increasing x . If $n>1$, y increases with increasing rapidity, if $n<1$, y increases with decreasing rapidity.

If the exponent is negative, the curves $y=ax^{-n}=\frac{a}{x^n}$ are hyperbolic curves, starting from $y=\infty$ for $x=0$, and decreasing to $y=0$ for $x=\infty$.

$n=1$ gives the straight line through the origin, $n=0$ and $n=\infty$ give, respectively, straight horizontal and vertical lines.

Figs. 61 to 71 give various curve shapes, corresponding to different values of n .

Parabolic Curves.

Fig. 61. $n=2$; $y=x^2$; the common parabola.

Fig. 62. $n=4$; $y=x^4$; the biquadratic parabola.

Fig. 63. $n=8$; $y=x^8$.

Fig. 64. $n=\frac{1}{2}$; $y=\sqrt{x}$; again the common parabola.

Fig. 65. $n=\frac{1}{4}$; $y=\sqrt[4]{x}$; the biquadratic parabola.

Fig. 66. $n=\frac{1}{8}$; $y=\sqrt[8]{x}$.

Hyperbolic Curves.

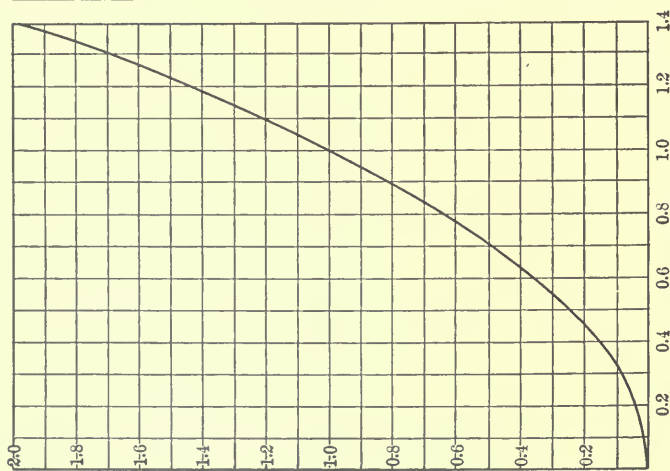
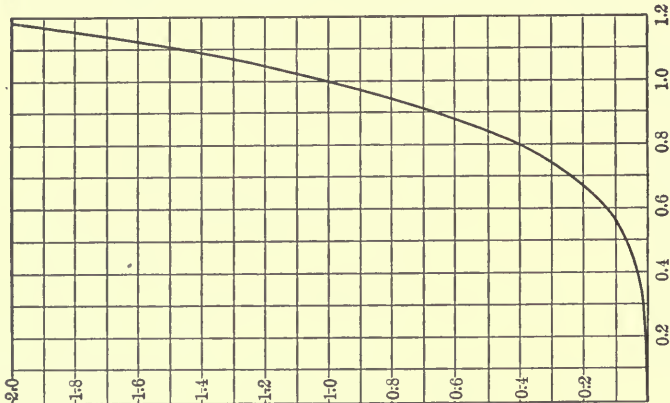
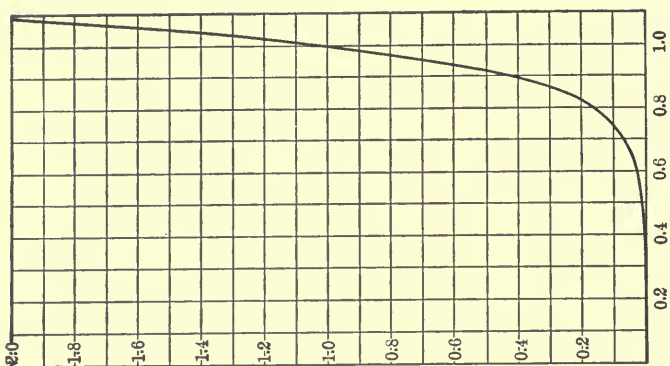
Fig. 67. $n=-1$; $y=\frac{1}{x}$; the equilateral hyperbola.

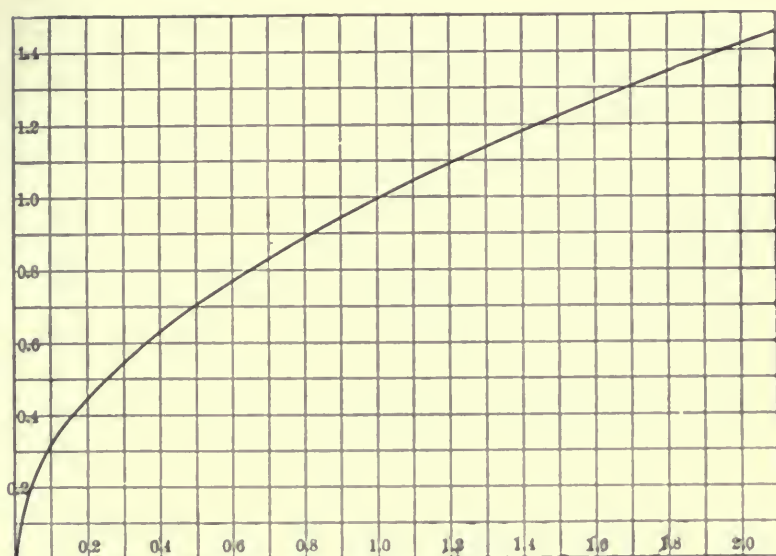
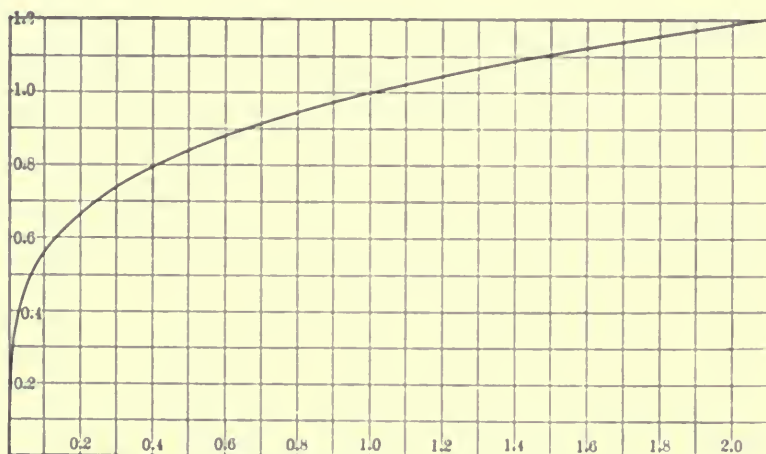
Fig. 68. $n=-2$; $y=\frac{1}{x^2}$.

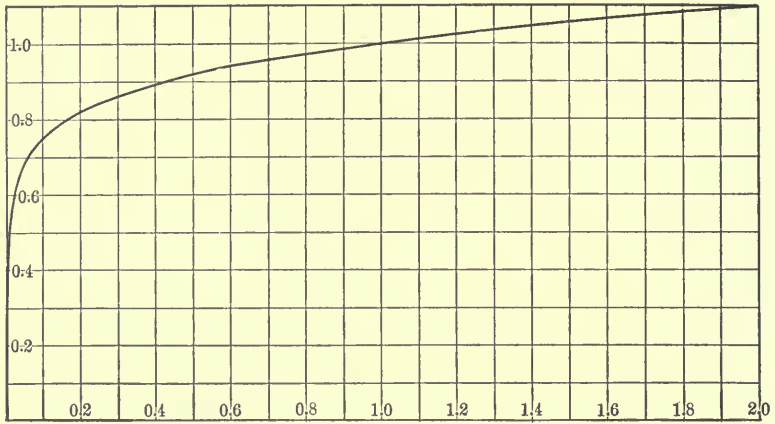
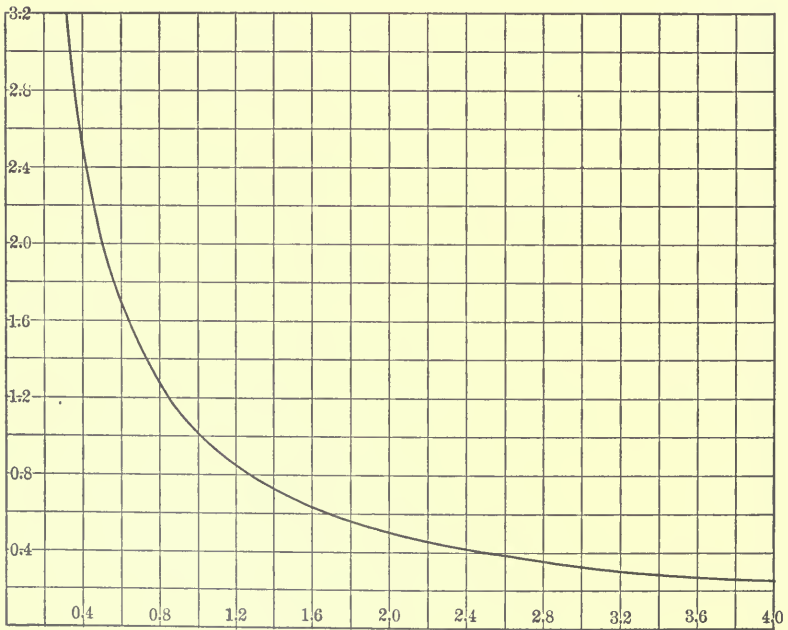
Fig. 69. $n=-4$; $y=\frac{1}{x^4}$.

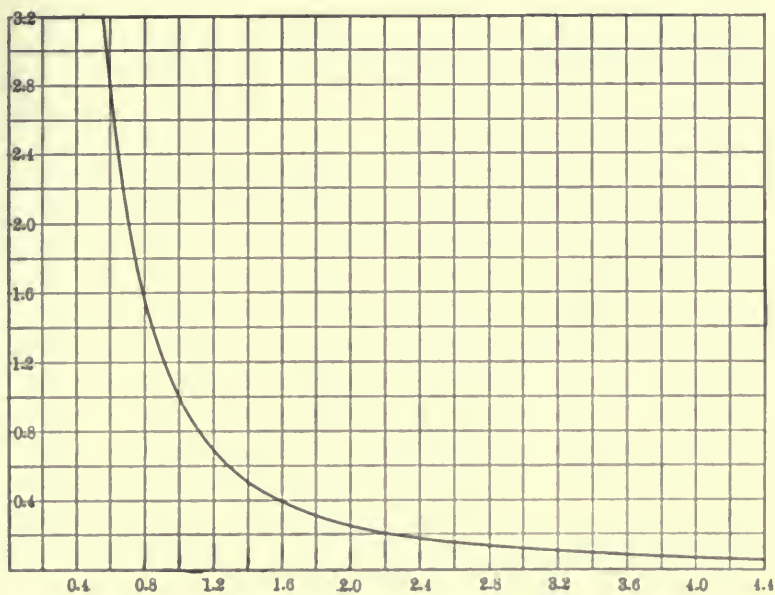
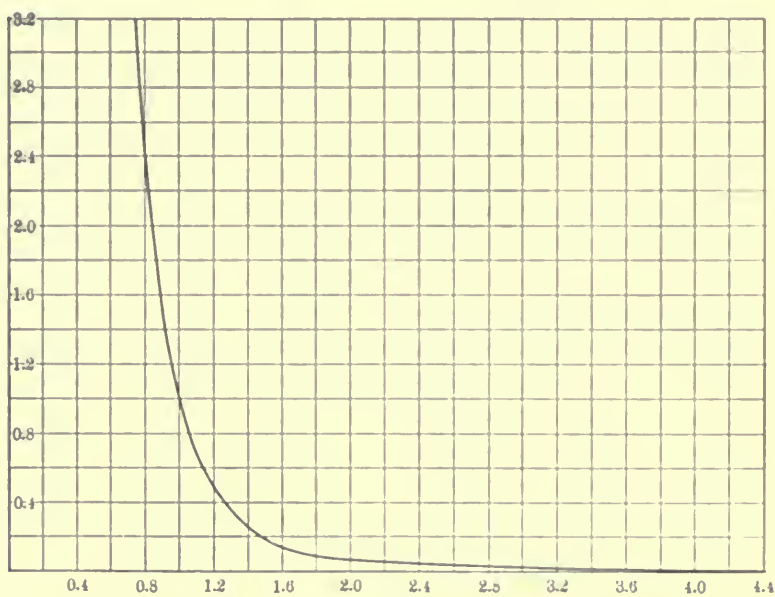
Fig. 70. $n=-\frac{1}{2}$; $y=\frac{1}{\sqrt{x}}$.

Fig. 71. $n=-\frac{1}{4}$; $y=\frac{1}{\sqrt[4]{x}}$.

FIG. 61. Parabolic Curve. $y = x^2$.FIG. 62. Parabolic Curve. $y = x^4$.FIG. 63. Parabolic Curve. $y = x^3$.

FIG. 64. Parabolic Curve. $y = \sqrt{x}$.FIG. 65. Parabolic Curve. $y = \sqrt[3]{x}$.

FIG. 66. Parabolic Curve. $y = \sqrt[6]{x}$.FIG. 67. Hyperbolic Curve (Equilateral Hyperbola). $y = \frac{1}{x}$.

FIG. 68. Hyperbolic Curve. $y = \frac{1}{x}$.FIG. 69. Hyperbolic Curve. $y = \frac{1}{x^4}$.

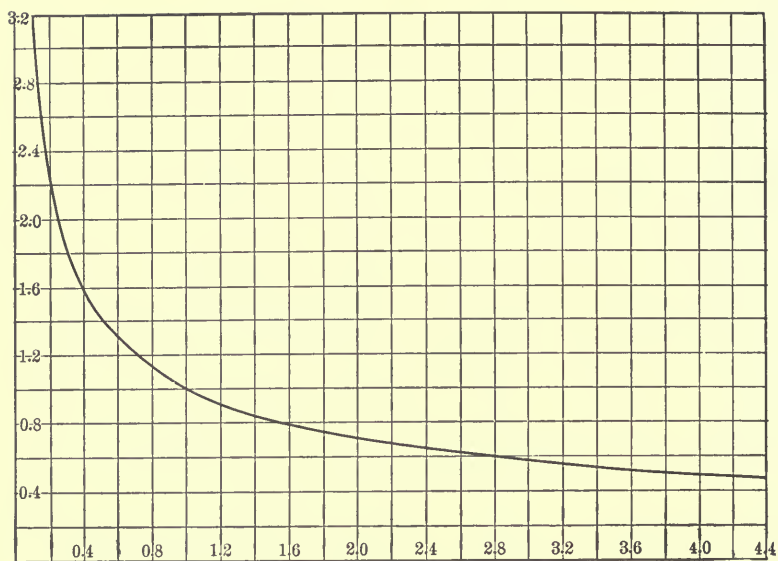


FIG. 70. Hyperbolic Curve. $y = \frac{1}{\sqrt{x}}$.

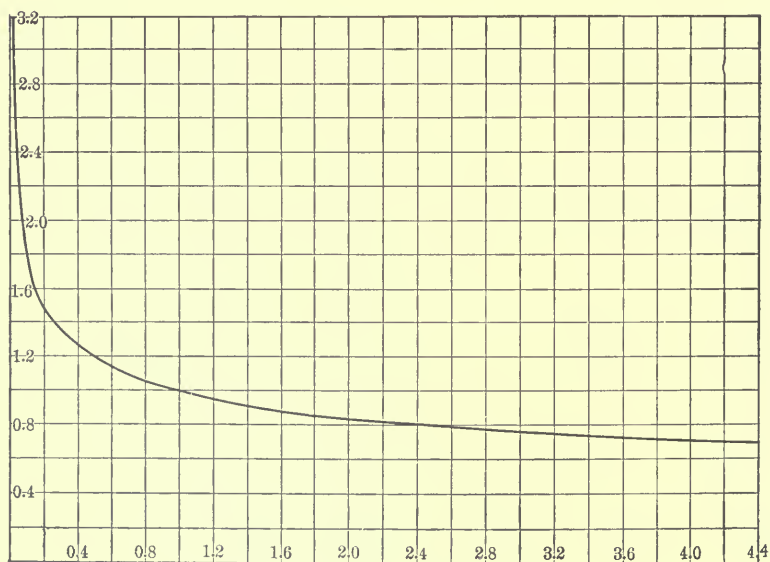


FIG. 71. Hyperbolic Curve. $y = \frac{1}{\sqrt[3]{x}}$.

In Fig. 72, sixteen different parabolic and hyperbolic curves are drawn together on the same sheet, for the following values: $n=1, 2, 4, 8, \infty; \frac{1}{2}, \frac{1}{4}, \frac{1}{8}, 0; -1, -2, -4, -8; -\frac{1}{2}, -\frac{1}{4}, -\frac{1}{8}$.

147. Parabolic and hyperbolic curves may easily be recognized by the fact that *if x is changed by a constant factor, y also changes by a constant factor.*

Thus, in the curve $y=x^2$, doubling the x increases the y fourfold; in the curve $y=x^{1.59}$, doubling the x increases the y threefold, etc.; that is, if in a curve,

$$y=f(x),$$

$$\frac{f(qx)}{f(x)} = \text{constant, for constant } q, \quad \dots \quad (17)$$

the curve is a parabolic or hyperbolic curve, $y=ax^n$, and

$$\frac{f(qx)}{f(x)} = \frac{a(qx)^n}{ax^n} = q^n. \quad \dots \quad (18)$$

If q is nearly 1, that is, the x is changed only by a small value, substituting $q=1+s$, where s is a small quantity, from equation (18),

$$\frac{f(x+sx)}{f(x)} = (1+s)^n = 1+ns;$$

hence,

$$\frac{f(x+sx)-f(x)}{f(x)} = ns; \quad \dots \quad (19)$$

that is, *changing x by a small percentage s , y changes by a proportional small percentage ns .*

Thus, parabolic and hyperbolic curves can be recognized by a small percentage change of x , giving a proportional small percentage change of y , and the proportionality factor is the exponent n ; or, they can be recognized by doubling x and seeing whether y hereby changes by a constant factor.

As illustration are shown in Fig. 73 the parabolic curves, which, for a doubling of x , increase y : 2, 3, 4, 5, 6, and 8 fold.

Unfortunately, this convenient way of recognizing parabolic and hyperbolic curves applies only if the curve passes through the origin, that is, has no constant term. If constant terms exist, as in equation (16), not x and y , but $(x-c)$ and $(y-b)$ follow the law of proportionate increases, and the recognition

becomes more difficult; that is, various values of c and of b are to be tried to find one which gives the proportionality.

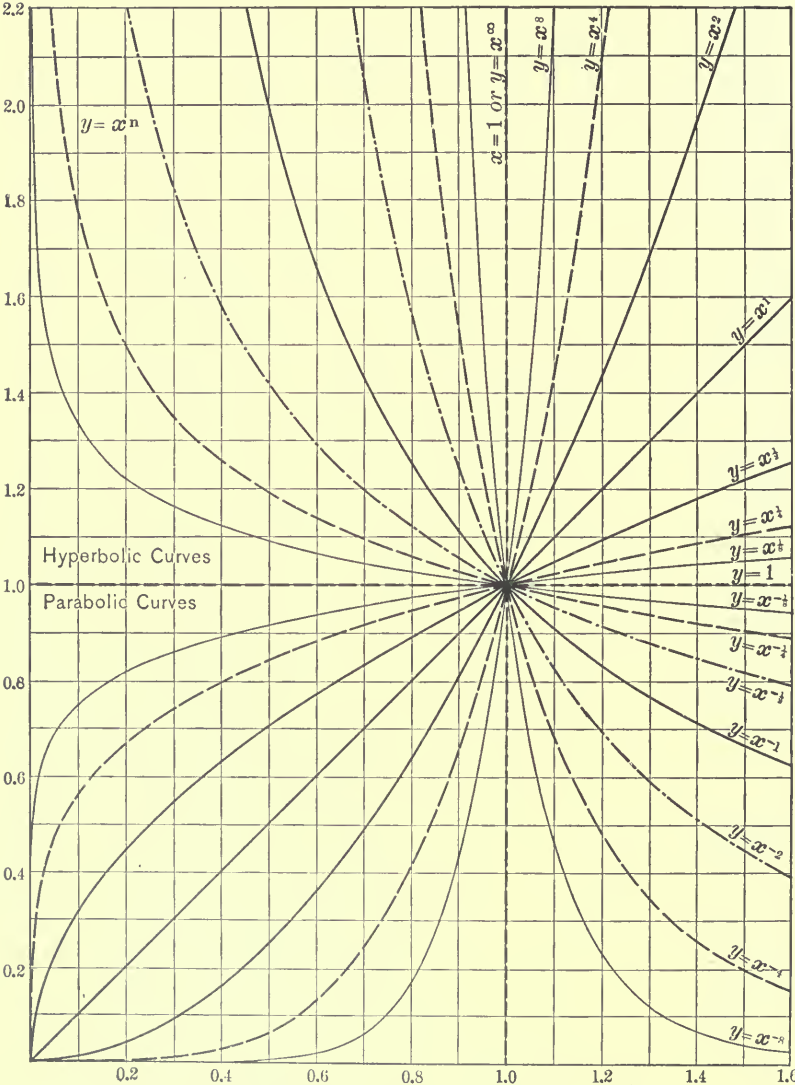


FIG. 72. Parabolic and Hyperbolic Curves. $y = x^n$.

148. Taking the logarithm of equation (15) gives
$$\log y = \log a + n \log x; \quad (20)$$

that is, a straight line; hence, a parabolic or hyperbolic curve can be recognized by plotting the logarithm of y against the logarithm of x . If this gives a straight line, the curve is parabolic or hyperbolic, and the slope of the logarithmic curve, $\tan \theta = n$, is the exponent.

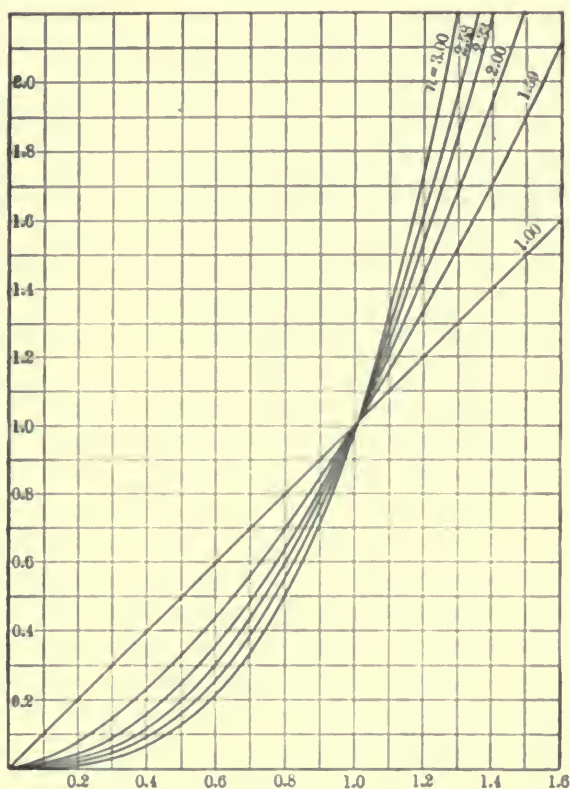


FIG. 73. Parabolic Curves. $y = x^n$.

This again applies only if the curve contain no constant term. If constant terms exist, the logarithmic line is curved. Therefore, by trying different constants c and b , the curvature of the logarithmic line changes, and by interpolation such constants can be found, which make the logarithmic line straight, and in this way, the constants c and b may be evaluated. If only one constant exist, that is, only b or only c , the process is relatively simple, but it becomes rather complicated with both

constants. This fact makes it all the more desirable to get from the physical nature of the problem some idea on the existence and the value of the constant terms.

Differentiating equation (20) gives:

$$\frac{dy}{y} = n \frac{dx}{x};$$

that is, in a parabolic or hyperbolic curve, the percentual change, or *variation* of y , is n times the percentual change, or variation of x , if n is the exponent.

Herefrom follows:

$$\frac{\frac{dy}{y}}{\frac{dx}{x}} = n.$$

that is, in a parabolic or hyperbolic curve, the *ratio of variation*,

$$m = \frac{\frac{dy}{y}}{\frac{dx}{x}}, \text{ is a constant, and equals the exponent } n.$$

Or, inversely:

If in an empirical curve the ratio of variation is constant the curve is—within the range, in which the ratio of variation is constant—a parabolic or hyperbolic curve, which has as exponent the ratio of variation.

In the range, however, in which the ratio of variation is not constant, it is *not* the exponent, and while the empirical curve might be expressed as a parabolic or hyperbolic curve with changing exponent (or changing coefficient), in this case the exponent may be very different from the ratio of variation, and the change of exponent frequently is very much smaller than the change of the ratio of variation.

This ratio of variation and exponent of the parabolic or hyperbolic approximation of an empirical curve must not be mistaken for each other, as has occasionally been done in reducing hysteresis curves, or radiation curves. They coincide

The curve, $y = e^{+x}$, has the same shape, except that the positive and the negative side (right and left) are interchanged.

Inverted these equations (21) to (24) may also be written thus,

$$\left. \begin{aligned} nx &= \log \frac{y}{a}; \\ n(x-c) &= \log \frac{y-b}{a}; \\ nx &= -\log \frac{y}{a}; \\ x &= -\log y; \end{aligned} \right\}, \dots \dots \dots (25)$$

that is, as logarithmic curves.

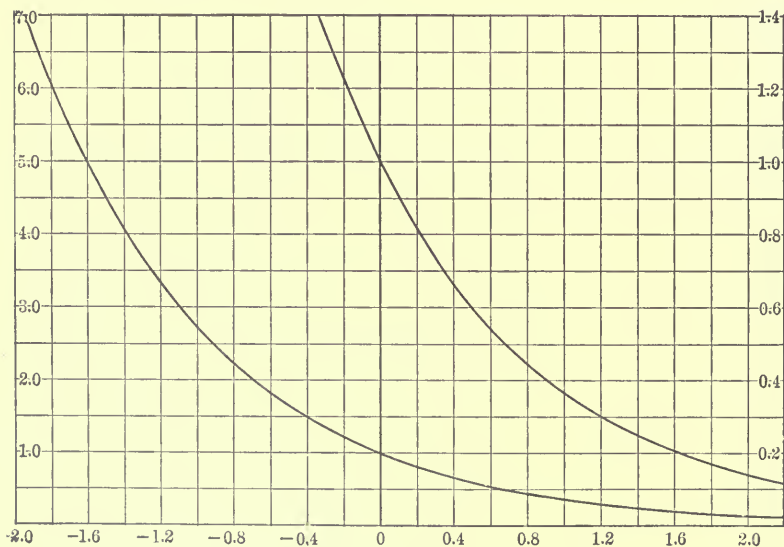


FIG. 74. Exponential Function. $y = e^{-x}$.

150. The characteristic of the exponential function (21) is, that an *increase of x by a constant term increases (or, in (23) and (24), decreases) y by a constant factor.*

Thus, if an empirical curve, $y = f(x)$, has such characteristic that

$$\frac{f(x+q)}{f(x)} = \text{constant, for constant } q, \dots \dots (26)$$

the curve is an exponential function, $y = a\epsilon^{nx}$, and the following equation may be written:

$$\frac{f(x+q)}{f(x)} = \frac{a\epsilon^{n(x+q)}}{a\epsilon^{nx}} = \epsilon^{nq}. \quad (27)$$

Hereby the exponential function can easily be recognized, and distinguished from the parabolic curve; in the former a constant *term*, in the latter a constant *factor* of x causes a change of y by a constant *factor*.

As result hereof, the exponential curve with negative exponent vanishes, that is, becomes negligibly small, with far greater rapidity than the hyperbolic curve, and the exponential

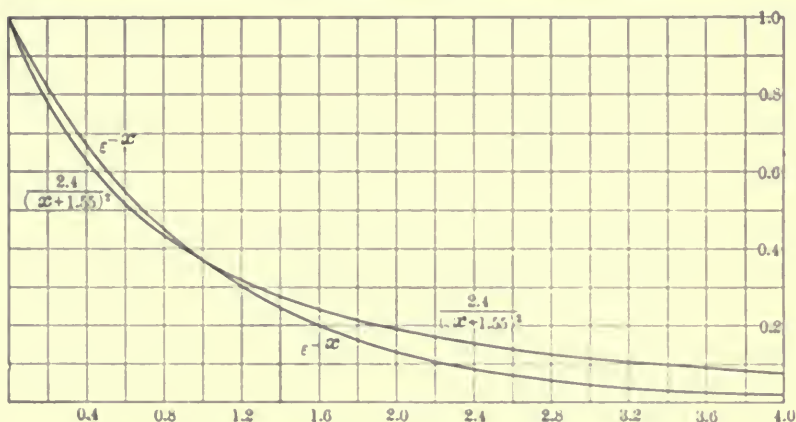


FIG. 75. Hyperbolic and Exponential Curves Comparison

function with positive exponent reaches practically infinite values far more rapidly than the parabolic curve. This is illustrated in Fig. 75, in which are shown superimposed the exponential curve, $y = \epsilon^{-x}$, and the hyperbolic curve, $y = \frac{2.4}{(x + 1.55)^2}$, which coincides with the exponential curve at $x = 0$ and at $x = 1$.

Taking the logarithm of equation (21) gives $\log y = \log a + nx \log \epsilon$, that is, $\log y$ is a linear function of x , and plotting $\log y$ against x gives a straight line. This is characteristic of

but would give a curve. In this case then $\log(y-b)$ would be plotted against x for various values of b , and by interpolation that value of b found which makes the logarithmic curve a straight line.

151. While the exponential function, when appearing singly, is easily recognized, this becomes more difficult with com-

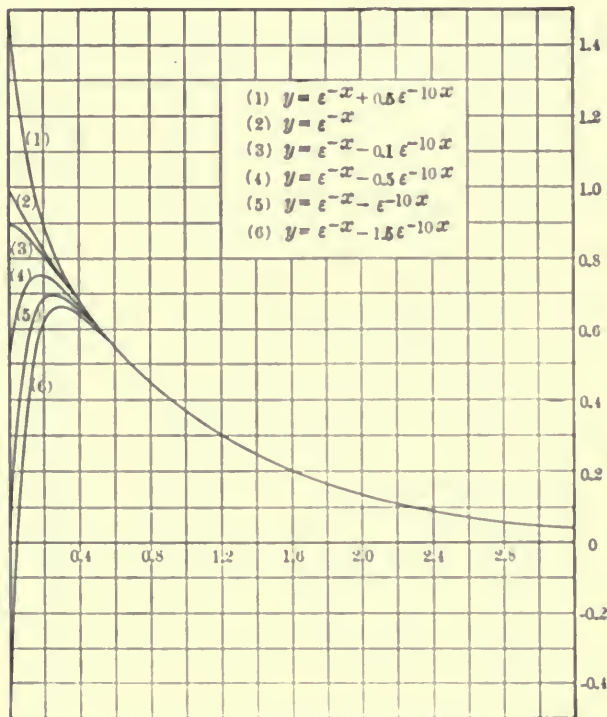


FIG. 77. Exponential Functions.

binations of two exponential functions of different coefficients in the exponent, thus,

$$y = a_1 e^{-c_1 x} \pm a_2 e^{-c_2 x}, \quad \dots \quad (29)$$

since for the various values of a_1 , a_2 , c_1 , c_2 , quite a number of various forms of the function appear.

As such a combination of two exponential functions frequently appears in engineering, some of the characteristic forms are plotted in Figs. 76 to 78.

Fig. 76 gives the following combinations of ε^{-x} and ε^{-2x} :

$$(1) \quad y = \varepsilon^{-x} + 0.5\varepsilon^{-2x};$$

$$(2) \quad y = \varepsilon^{-x} + 0.2\varepsilon^{-2x};$$

$$(3) \quad y = \varepsilon^{-x};$$

$$(4) \quad y = \varepsilon^{-x} - 0.2\varepsilon^{-2x};$$

$$(5) \quad y = \varepsilon^{-x} - 0.5\varepsilon^{-2x};$$

$$(6) \quad y = \varepsilon^{-x} - 0.8\varepsilon^{-2x};$$

$$(7) \quad y = \varepsilon^{-x} - \varepsilon^{-2x};$$

$$(8) \quad y = \varepsilon^{-x} - 1.5\varepsilon^{-2x}.$$

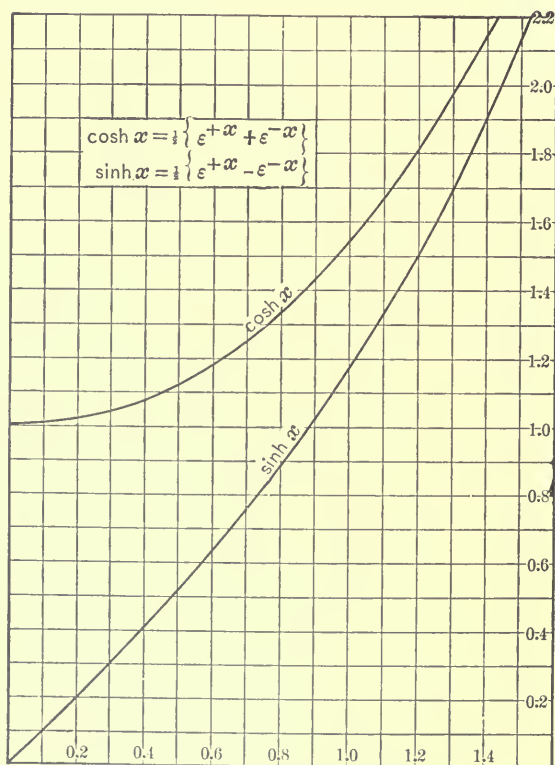


FIG. 78. Hyperbolic Functions.

Fig. 77 gives the following combination of ϵ^{-x} and ϵ^{-10x}

$$(1) \quad y = \epsilon^{-x} + 0.5\epsilon^{-10x};$$

$$(2) \quad y = \epsilon^{-x};$$

$$(3) \quad y = \epsilon^{-x} - 0.1\epsilon^{-10x};$$

$$(4) \quad y = \epsilon^{-x} - 0.5\epsilon^{-10x};$$

$$(5) \quad y = \epsilon^{-x} - \epsilon^{-10x};$$

$$(6) \quad y = \epsilon^{-x} - 1.5\epsilon^{-10x}.$$

Fig. 78 gives the hyperbolic functions as combinations of ϵ^{+x} and ϵ^{-x} thus,

$$y = \cosh x = \frac{1}{2}(\epsilon^{+x} + \epsilon^{-x});$$

$$y = \sinh x = \frac{1}{2}(\epsilon^{+x} - \epsilon^{-x}).$$

C. Evaluation of Empirical Curves.

152. In attempting to solve the problem of finding a mathematical equation, $y=f(x)$, for a series of observations or tests, the corresponding values of x and y are first tabulated and plotted as a curve.

From the nature of the physical problem, which is represented by the numerical values, there are derived as many data as possible concerning the nature of the curve and of the function which represents it, especially at the zero values and the values at infinity. Frequently hereby the existence or absence of constant terms in the equation is indicated.

The $\log x$ and $\log y$ are tabulated and curves plotted between x , y , $\log x$, $\log y$, and seen, whether some of these curves is a straight line and thereby indicates the exponential function, or the parabolic or hyperbolic function.

If cross-section paper is available, having both coordinates divided in logarithmic scale, and also cross-section paper having one coordinate divided in logarithmic, the other in common scale, x and y can be directly plotted on these two forms of logarithmic cross-section paper. Usually not much is saved thereby, as for the numerical calculation of the constants the logarithms still have to be tabulated.

If neither of the four curves: $x, y; x, \log y; \log x, y; \log x, \log y$ is a straight line, and from the physical condition the absence of a constant term is assured, the function is neither an exponential nor a parabolic or hyperbolic. If a constant term is probable or possible, curves are plotted between $x, y-b, \log x, \log (y-b)$ for various values of b , and if hereby one of the curves straightens out, then, by interpolation, that value of b is found, which makes one of the curves a straight line, and thereby gives the curve law. A convenient way of doing this is: if the curve with $\log y$ (curve O) is curved by angle α_0 (α_0 being for instance the angle between the tangents at the two end points of the curve, or the difference of the slopes at the two end points), use a value b_1 , and plot the curve with $\log (y-b_1)$ (curve 1), and observe its curvature α_1 . Then interpolate a value b_2 , between b_1 and O , in proportion to the curvatures α_1 and α_0 , and plot curve with $\log (y-b_2)$ (curve 2), and again interpolate a value b_3 between b_2 and either b_1 or O , whichever curve is nearer in slope to curve 2, continue until either the curve with $\log (y-b)$ becomes a straight line, or an S curve and in this latter case shows that the empirical curve cannot be represented in this manner.

In this work, logarithmic paper is very useful, as it permits plotting the curves without first looking up the logarithms, the latter being done only when the last approximation of b is found. In the same manner, if a constant term is suspected in the x , the value $(x-c)$ is used and curves plotted for various values of c . Frequently the existence and the character of a constant term is indicated by the shape of the curve; for instance, if one of the curves plotted between $x, y, \log x, \log y$ approaches straightness for high, or for low values of the abscissas, but curves considerably at the other end, a constant term may be suspected, which becomes less appreciable at one end of the range. For instance, the effect of the constant c in $(x-c)$ decreases with increase of x .

Sometimes one of the curves may be a straight line at one end, but curve at the other end. This may indicate the presence of a term, which vanishes for a part of the observations. In this case only the observations of the range which gives a straight line are used for deriving the curve law, the curve calculated therefrom, and then the difference between the calculated curve and the observations further investigated.

Such a deviation of the curve from a straight line may also indicate a change of the curve law, by the appearance of secondary phenomena, as magnetic saturation, and in this case, an equation may exist only for that part of the curve where the secondary phenomena are not yet appreciable. The same equation may then be applied to the remaining part of the curve, by assuming one of the constants, as a coefficient, or an exponent, to change. Or a second equation may be derived for this part of the curve and one part of the curve represented by one, the other by another equation. The two equations may then overlap, and at some point the curve represented equally well by either equation, or the ranges of application of the two equations may be separated by a transition range, in which neither applies exactly.

If neither the exponential functions nor the parabolic and hyperbolic curves satisfactorily represent the observations, further trials may be made by calculating and tabulating $\frac{x}{y}$

and $\frac{y}{x}$, and plotting curves between x , y , $\frac{x}{y}$, $\frac{y}{x}$. Also expressions as $x^2 + by^2$, and $(x-a)^2 + b(y-c)^2$, etc., may be studied.

Theoretical reasoning based on the nature of the phenomenon represented by the numerical data frequently gives an indication of the form of the equation, which is to be expected, and inversely, after a mathematical equation has been derived a trial may be made to relate the equation to known laws and thereby reduce it to a rational equation.

In general, the resolution of empirical data into a mathematical expression largely depends on trial, directed by judgment based on the shape of the curve and on a knowledge of the curve shapes of various functions, and only general rules can thus be given.

A number of examples may illustrate the general methods of reduction of empirical data into mathematical functions.

153. Example 1. In a 118-volt tungsten filament incandescent lamp, corresponding values of the terminal voltage e and the current i are observed, that is, the so-called "volt-ampere characteristic" is taken, and therefrom an equation for the volt-ampere characteristic is to be found.

The corresponding values of e and i are tabulated in the first two columns of Table III and plotted as curve I in

Fig. 79. In the third and fourth column of Table III are given $\log e$ and $\log i$. In Fig. 79 then are plotted $\log e, i$, as curve II; $e, \log i$, as curve III; $\log e, \log i$, as curve IV.

As seen from Fig. 79, curve IV is a straight line, that is

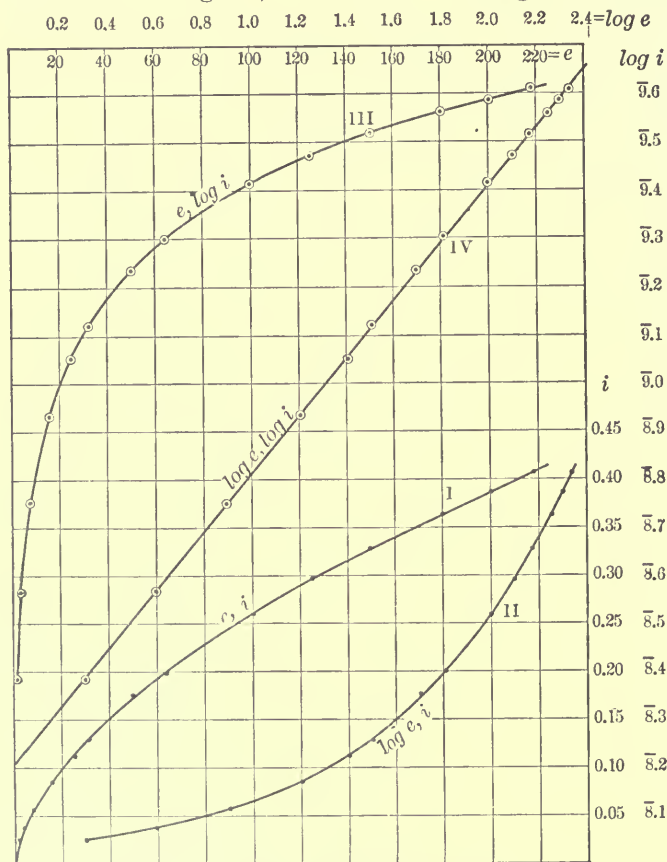


Fig. 79. Investigation of Volt-ampere Characteristic of Tungsten Lamp Filament.

$$\log i = A + n \log e; \quad \text{or} \quad i = ae^n,$$

which is a parabolic curve.

The constants a and n may now be calculated from the numerical data of Table III by the method of least squares, as discussed in Chapter IV, paragraph 120. While this method gives the most accurate results, it is so laborious as to be seldom used in engineering; generally, values of the constants a and n , sufficiently accurate for most practical

purposes, are derived by the so-called " $\Sigma\Delta$ method," which, with proper tabular arrangement of the numerical values, gives high accuracy with a minimum of work.

TABLE III.

VOLT-AMPERE CHARACTERISTIC OF 118-VOLT TUNGSTEN LAMP.

e	i	$\log e$	$\log i$	$\bar{8}.211 + 0.6 \log e$	J
2	0.0245	0.301	$\bar{5}.392$	$\bar{5}.389$	-0.003
4	0.037	0.802	$\bar{5}.568$	$\bar{5}.572$	-0.004
8	0.0568	0.903	$\bar{5}.754$	$\bar{5}.753$	+0.001
16	0.0855	1.204	$\bar{5}.932$	$\bar{5}.933$	-0.001
25	0.1125	1.398	$\bar{5}.051$	$\bar{5}.050$	+0.001
32	0.1295	1.505	$\bar{5}.112$	$\bar{5}.114$	-0.002
50	0.1715	1.699	$\bar{5}.234$	$\bar{5}.230$	+0.004
64	0.200	1.806	$\bar{5}.301$	$\bar{5}.295$	+0.006
100	0.2605	2.000	$\bar{5}.418$	$\bar{5}.411$	+0.005
125	0.2965	2.097	$\bar{5}.472$	$\bar{5}.469$	+0.003
150	0.3295	2.176	$\bar{5}.518$	$\bar{5}.518$	0
180	0.3635	2.255	$\bar{5}.561$	$\bar{5}.564$	-0.003
200	0.3865	2.301	$\bar{5}.587$	$\bar{5}.592$	-0.005
218	0.407	2.338	$\bar{5}.610$	$\bar{5}.614$	-0.004

$\Sigma 7 = 7.612$	$\Sigma 0.43$	avg ± 0.003
$\Sigma 7 = 14.973$	$\bar{8}.465$	-4.7 per cent

$J = 7.361$	4.422	
$n = \frac{4.422}{7.361} =$	0.6007 ≈ 0.6	
$\Sigma 14 = 22.585$	$\bar{8}.505$	
$0.6 \times 22.585 =$	13.551	
	$J = \bar{8}.505 - 13.551 = \bar{4}.954$	
$\bar{4}.954 + 14 =$	$\bar{8}.211$	

$\log i = \bar{8}.211 + 0.6 \log e$ and $i = 0.01625e^{0.6}$

The fourteen sets of observations are divided into two groups of seven each, and the sums of $\log e$ and $\log i$ formed. They are indicated as $\Sigma 7$ in Table III.

Then subtracting the two groups $\Sigma 7$ from each other, eliminates A , and dividing the two differences J , gives the exponent, $n=0.6011$; this is so near to 0.6 that it is reasonable to assume that $n=0.6$, and this value then is used.

Now the sum of all the values of $\log e$ is formed, given as $\Sigma 14$ in Table III, and multiplied with $n=0.6$, and the product

subtracted from the sum of all the $\log i$. The difference A then equals $14A$, and, divided by 14, gives

$$A = \log a = \bar{8}.211;$$

hence, $a = 0.01625$, and the volt-ampere characteristic of this tungsten lamp thus follows the equation,

$$\begin{aligned}\log i &= \bar{8}.211 + 0.6 \log e; \\ i &= 0.01625e^{0.6}.\end{aligned}$$

From e and i can be derived the power input $p = ei$, and the resistance $r = \frac{e}{i}$:

$$p = 0.01625e^{1.6};$$

$$r = \frac{e^{0.4}}{0.01625},$$

and, eliminating e from these two expressions, gives

$$p = 0.01625^5 r^4 = 11.35 r^4 \times 10^{-10},$$

that is, the power input varies with the fourth power of the resistance.

Assuming the resistance r as proportional to the absolute temperature T , and considering that the power input into the lamp is radiated from it, that is, is the power of radiation P_r , the equation between p and r also is the equation between P_r and T , thus,

$$P_r = kT^4;$$

that is, the radiation is proportional to the fourth power of the absolute temperature. This is the law of black body radiation, and above equation of the volt-ampere characteristic of the tungsten lamp thus appears as a conclusion from the radiation law, that is, as a rational equation.

154. Example 2. In a magnetite arc, at constant arc length, the voltage consumed by the arc, e , is observed for different values of current i . To find the equation of the volt-ampere characteristic of the magnetite arc:

TABLE IV.
VOLT-AMPERE CHARACTERISTIC OF MAGNETITE ARC.

i	e	$\log i$	$\log e$	$(e-40)$	$\log (e-40)$	$(e-30)$	$\log (e-30)$	e_c	J
0.5	180	9.699	2.204	120	2.079	130	2.114	158	-2
1	120	0.000	2.079	80	1.903	90	1.954	120.4	+0.4
2	94	0.301	1.973	54	1.732	64	1.808	94	0
4	75	0.602	1.875	35	1.544	45	1.653	75.2	+0.2
8	62	0.903	1.792	22	1.342	32	1.505	62	0
12	56	1.079	1.748	16	1.204	26	1.415	56.2	+0.2

$\Sigma 3 = 0.000$							5.674
$\Sigma 3 = 2.584$							4.573
$J = 2.584$							1.301
				$n = \frac{-1.301}{2.584} = -0.5034 \approx -0.5$			
$\Sigma 6 = 2.584$							10.447
			2.584×0.5				-1.292
						$J = 11.739$	
						$11.739 + 0 = 1.956 =$	
						$\log (e-30) = 1.956 = 0.5 \log i$	
						$e-30 = 90.41 \approx 90$ and $e-30 = \frac{50}{\sqrt{1}}$	

The first four columns of Table IV give i , e , $\log i$, $\log e$. Fig. 80 gives the curves: i , e , as I; i , $\log e$, as II; $\log i$, e , as III; $\log i$, $\log e$, as IV.

Neither of these curves is a straight line. Curve IV is relatively the straightest, especially for high values of e . This points toward the existence of a constant term. The existence of a constant term in the arc voltage, the so-called "counter e.m.f. of the arc" is physically probable. In Table IV thus are given the values $(e-40)$ and $\log (e-40)$, and plotted as curve V. This shows the opposite curvature of IV. Thus the constant term is less than 40. Estimating by interpolation, and calculating in Table IV $(e-30)$ and $\log (e-30)$, the latter, plotted against $\log i$ gives the straight line VI. The curve law thus is

$$\log (e-30) = A + n \log i.$$

Proceeding in Table IV in the same manner with $\log i$ and $\log (e-30)$ as was done in Table III with $\log e$ and $\log i$, gives

$$n = -0.5; \quad A = \log a = 1.956; \quad \text{and} \quad a = 90.4;$$

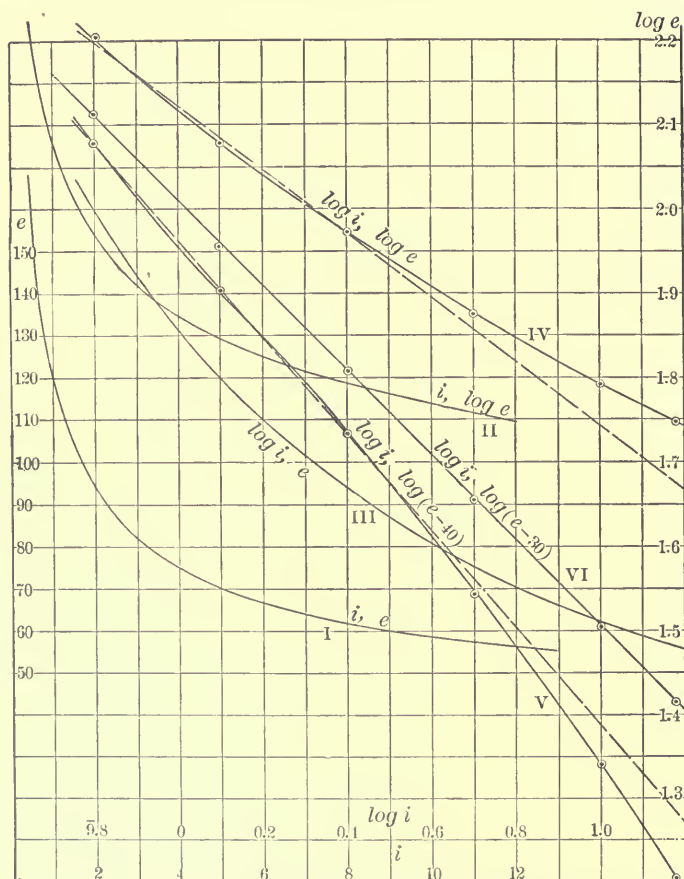


FIG. 80. Investigation of Volt-ampere Characteristic of Magnetite Arc.

hence

$$\log (e-30) = 1.956 - 0.5 \log i;$$

$$e-30 = 90.4 i^{-0.5};$$

$$e = 30 + \frac{90.4}{\sqrt{i}}$$

which is the equation of the magnetite arc volt-ampere characteristic.

155. Example 3. The change of current resulting from a change of the conditions of an electric circuit containing resistance, inductance, and capacity is recorded by oscillograph and gives the curve reproduced as I in Fig. 81. From this curve

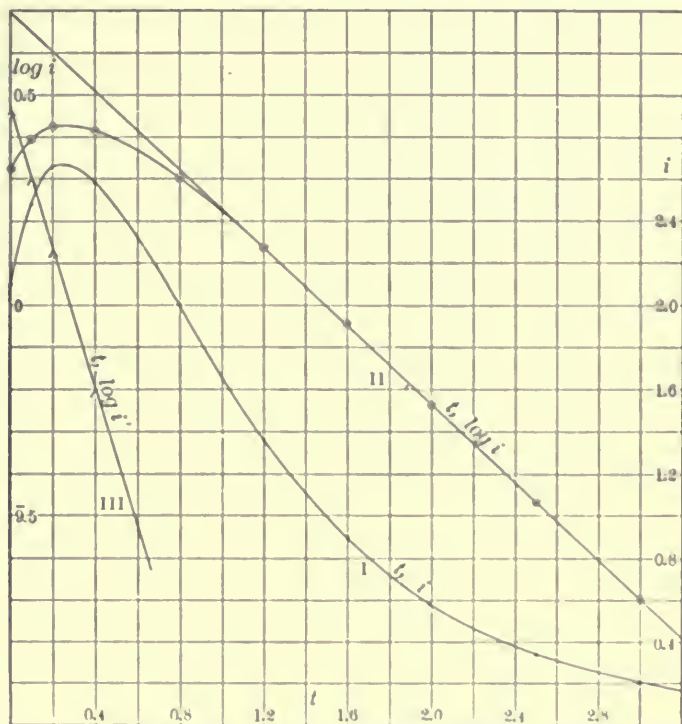


FIG. 81. Investigation of Curve of Current Change in Electric Circuit.

are taken the numerical values tabulated as t and i in the first two columns of Table V. In the third and fourth columns are given $\log t$ and $\log i$, and curves then plotted in the usual manner. Of these curves only the one between t and $\log i$ is shown, as II in Fig. 81, since it gives a straight line for the higher values of t . For the higher values of t , therefore,

$$\log i = A - nt; \quad \text{or,} \quad i = a e^{-nt};$$

that is, it is an exponential function.

TABLE V.
TRANSIENT CURRENT CHARACTERISTICS.

	t	i	$\log t$	$\log i$	i_1	i'	t	$\log i'$	i_2	i_c	A
	0	2.10	—	0.322	4.94	2.84	0	0.461	2.85	2.09	-0.01
	0.1	2.48	5.000	0.394	4.44	1.96	0.1	0.292	1.94	2.50	+0.02
	0.2	2.66	5.301	0.425	3.98	1.32	0.2	0.121	1.32	2.66	0
	0.4	2.58	5.602	0.412	3.21	0.63	0.4	5.799	0.61	2.60	+0.02
	0.8	2.00	5.903	0.301	2.09	0.09	0.8	8.954	0.13	1.96	-0.04
	1.2	1.36	0.079	0.134	1.36	0	—	—	0.03	1.33	-0.03
	1.6	0.90	0.204	5.954	0.89	-0.01	—	—	0.01	0.88	-0.02
	2.0	0.58	0.301	5.763	0.58	0	—	—	—	0.58	0
	2.5	0.34	0.398	5.531	0.34	0	—	—	—	0.34	0
	3.0	0.20	0.477	5.301	0.20	0	—	—	—	0.20	0
$\begin{aligned} x_3 &= \frac{4.8}{3} = 1.6 \\ x_2 &= \frac{5.5}{2} = 2.75 \\ A &= 1.15 \\ \log \epsilon \times 1.15 &= 0.499 \\ n_1 &= \frac{-0.534}{0.499} = -1.07 \end{aligned}$						$\begin{aligned} x_2 &= -0.1 \quad 0.753 \\ x_2 &= -0.8 \quad 9.920 \\ A &= 0.5 - 0.833 \\ \log \epsilon \times 0.5 &= 0.217 \\ n_2 &= \frac{-0.833}{0.217} = -3.84 \\ x_4 &= 0.7 \quad 0.673 \\ n_2 \log \epsilon \times 0.7 &= -1.167 \\ A &= 1.840 \\ 1.840 \div 4 &= 0.460 = A_2 = \log \alpha_2 \\ \alpha_2 &= 2.85 \\ \log i_2 &= 0.460 - 3.84t \log \epsilon \\ i_2 &= 2.85\epsilon^{-3.84t} \end{aligned}$					
$\begin{aligned} x_3 &= 10.3 \\ 10.3 \times n_1 \log \epsilon &= -4.784 \\ A &= 3.467 \\ 3.467 \div 5 &= 0.693 = A_1 = \log \alpha_1 \\ \alpha_1 &= 4.94 \\ \log i_1 &= 0.693 - 1.07t \log \epsilon \\ i_1 &= 4.94\epsilon^{-1.07t} \end{aligned}$						$\begin{aligned} i_c &= 4.94\epsilon^{-1.07t} - 2.85\epsilon^{-3.84t} \end{aligned}$					

To calculate the constants a and n , the range of values is used, in which the curve II is straight; that is, from $t=1.2$ to $t=3$. As these are five observations, they are grouped in two pairs, the first 3, and the last 2, and then for t and $\log i$, one-third of the sum of the first 3, and one-half of the sum of the last 2 are taken. Subtracting, this gives,

$$At = 1.15; \quad A \log i = -0.534.$$

Since, however, the equation, $i = a\epsilon^{-nt}$, when logarithmated, gives

$$\log i = \log a - nt \log \epsilon,$$

thus

$$A \log i = -n \log \epsilon At,$$

it is necessary to multiply Jt by $\log \epsilon = 0.4343$ before dividing it into $\log i$ to derive the value of n . This gives $n = 1.07$.

Taking now the sum of all the five values of t , multiplying it by $\log \epsilon$, and subtracting this from the sum of all the five values of $\log i$, gives $5A = 3.467$; hence

$$A = \log a = 0.693,$$

$$a = 4.94,$$

and

$$\log i_1 = 0.693 - 1.07t \log \epsilon;$$

$$i_1 = 4.94\epsilon^{-1.07t}.$$

The current i_1 is calculated and given in the fifth column of Table V, and the difference $i' = J = i_1 - i$ in the sixth column. As seen, from $t = 1.2$ upward, i_1 agrees with the observations. Below $t = 1.2$, however, a difference i' remains, and becomes considerable for low values of t . This difference apparently is due to a second term, which vanishes for higher values of t . Thus, the same method is now applied to the term i' ; column 8 gives $\log i'$, and in curve III of Fig. 81 is plotted $\log i'$ against t . This curve is seen to be a straight line, that is, i' is an exponential function of t .

Resolving i' in the same manner, by using the first four points of the curve, from $t = 0$ to $t = 0.4$, gives

$$\log i_2 = 0.460 - 3.84t \log \epsilon;$$

$$i_2 = 2.85\epsilon^{-3.84t},$$

and, therefore,

$$i = i_1 - i_2 = 4.94\epsilon^{-1.07t} - 2.85\epsilon^{-3.84t}$$

is the equation representing the current change.

The numerical values are calculated from this equation and given under i_c in Table V, the amount of their difference from the observed values are given in the last column of this table.

A still greater approximation may be secured by adding the calculated values of i_2 to the observed values of i in the last five observations, and from the result derive a second approximation of i_1 , and by means of this a second approximation of i_2 .

156. As further example may be considered the resolution of the core loss curve of an electric motor, which had been expressed irrationally by a potential series in paragraph 144 and Table I.

TABLE VI.
CORE LOSS CURVE.

e Volts.	P_i kw.	$\log e$	$\log P_i$	$1.6 \log e$	$A = \log P_i$ $- 1.6 \log e$	P_c	$\mathcal{A}$
40	0.63	1.602	0.799	2.563	7.238	0.70	-0.07
60	1.36	1.778	0.134	2.845	7.289	1.34	+0.02
80	2.18	1.903	0.338	3.045	7.293	2.12	+0.06
100	3.00	2.000	0.477	3.200	7.277	3.03	-0.03
120	3.93	2.079	0.594	3.326	7.268	4.05	-0.12
140	6.22	2.146	0.794	3.434	7.360	5.20	+1.02
160	8.59	2.204	0.934	3.526	7.408	6.43	+2.16
$\Sigma 3 = 5.283$ 0.271 $\log P_i = 7.282 + 1.6 \log e$ $\Sigma 3 \div 3 = 1.761$ 0.090 $P_i = 1.914e^{1.6}$, in watts $\Sigma 2 = 4.079$ 1.071 $\Sigma 2 \div 2 = 2.0395$ 0.535 $\mathcal{A} = 0.2785$ 0.445 $n = \frac{0.445}{0.2785} = 1.598 \approx 1.6$							

The first two columns of Table VI give the observed values of the voltage e and the core loss P_i in kilowatts. The next two columns give $\log e$ and $\log P_i$. Plotting the curves shows that $\log e$, $\log P_i$ is approximately a straight line, as seen in Fig. 82, with the exception of the two highest points of the curve.

Excluding therefore the last two points, the first five observations give a parabolic curve.

The exponent of this curve is found by Table VI as $n=1.598$; that is, with sufficient approximation, as $n=1.6$.

To see how far the observations agree with the curve, as given by the equation,

$$P_i = ae^{1.6}$$

in the fifth column $1.6 \log e$ is recorded, and in the sixth column, $A = \log a = \log P_i - 1.6 \log e$. As seen, the first and the last two values of A differ from the rest. The first value corre-

sponds to such a low value of P_i as to lower the accuracy of the observation. Averaging then the four middle values, gives $A = \bar{7}.282$; hence,

$$\log P_i = \bar{7}.282 + 1.6 \log e,$$

$$P_i = 1.914e^{1.6}, \text{ in watts.}$$

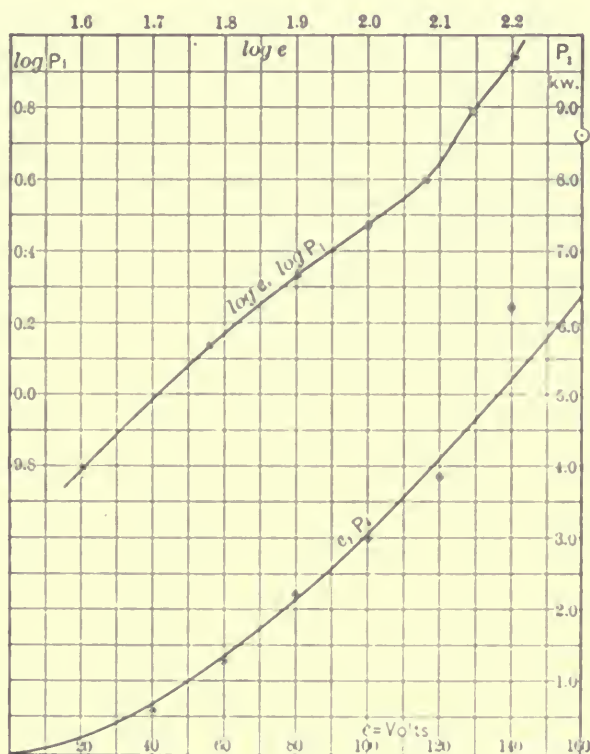


FIG. 82. Investigation of Curves.

This equation is calculated, as P_c , and plotted in Fig. 82. The observed values of P_i are marked by circles. As seen, the agreement is satisfactory, with the exception of the two highest values, at which apparently an additional loss appears, which does not exist at lower voltages. This loss probably is due to eddy currents caused by the increasing magnetic stray field resulting from magnetic saturation.

Excluding the three lowest values of the observations, as not lying on the straight line, from the remaining eight values, as calculated in Table VII, the following relation is derived,

$$\frac{H}{B} = 0.211 + 0.0507 H,$$

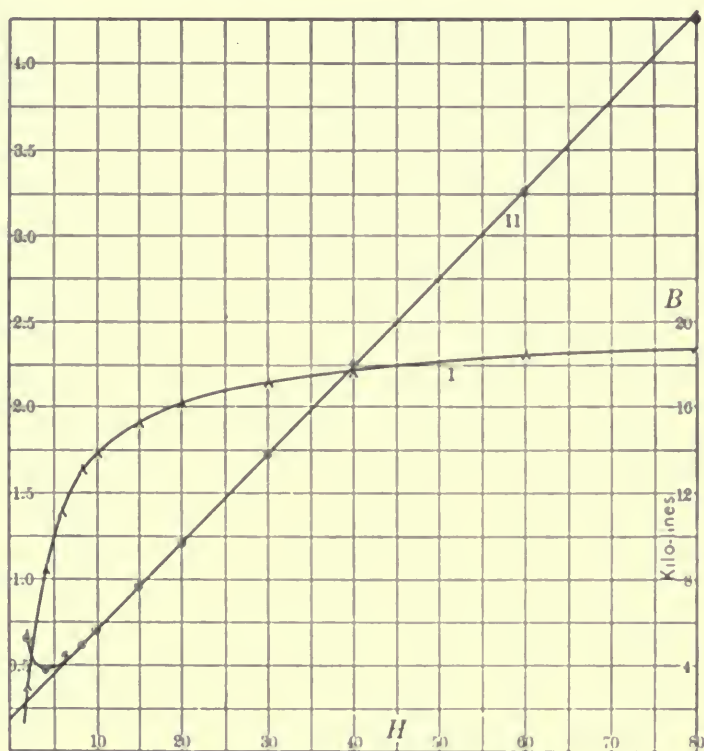


FIG. S3. Investigation of Magnetization Curve.

and herefrom,

$$B = \frac{H}{0.211 + 0.0507 H}$$

is the equation of the magnetic characteristic for values of H from eight upward.

The values calculated from this equation are given as B_c in Table VII.

The difference between the observed values of $\frac{H}{B}$, and the value given by above equation, which is appreciable up to $H=6$, could now be further investigated, and would be found to approximately follow an exponential law.

As a final example may be considered the investigation of a hysteresis curve of silicon steel, of which the numerical values are given in columns 1 and 2 of Table VIII.

The first column gives the magnetic density B , in lines of magnetic force per cm.²; the second column the hysteresis loss w , in ergs per cycle per kg. (specific density 7.5). The third column gives $\log B$, and the fourth column $\log w$.

Of the four curves between B , w , $\log B$, $\log w$, only the curve relating $\log w$ to $\log B$ approximates a straight line, and is given in the upper part of Fig. 84. This curve is not a straight line throughout its entire length, but only two sections of it are straight, from $B=50$ to $B=400$, and from $B=1600$ to $B=8000$, but the curve bends between 500 and 1200, and above 8000.

Thus two empirical formulas, of the form: $w = aB^n$, are calculated, in the usual manner, in Table VIII. The one applies for lower densities, the other for medium densities:

$$\text{Low density:} \quad B \leq 400: \quad w = 0.00341B^{2.11}$$

$$\text{Medium density: } 1600 \leq B \leq 8000: \quad w = 0.1096B^{1.60}$$

In Table VIII the values for the lower range are denoted by the index 1, for the higher range by the index 2.

Neither of these empirical formulas applies strictly to the range: $400 < B < 1600$, and to the range $B > 8000$. They may be applied within these ranges, by assuming either the coefficient a as varying, or the exponent n as varying, that is, applying a correction factor to a , or to n .

Thus, in the range: $400 < B < 1600$, the loss may be represented by:

(1) An extension of the low density formula:

$$w = a_1 B^{2.11} \quad \text{or} \quad w = 0.00341 B^{n_1}.$$

(2) An extension of the medium density formula

$$w = a_2 B^{1.6} \quad \text{or} \quad w = 0.1096 B^{n_2},$$

by giving tables or curves of a respectively n . Such tables are most conveniently given as a percentage correction.

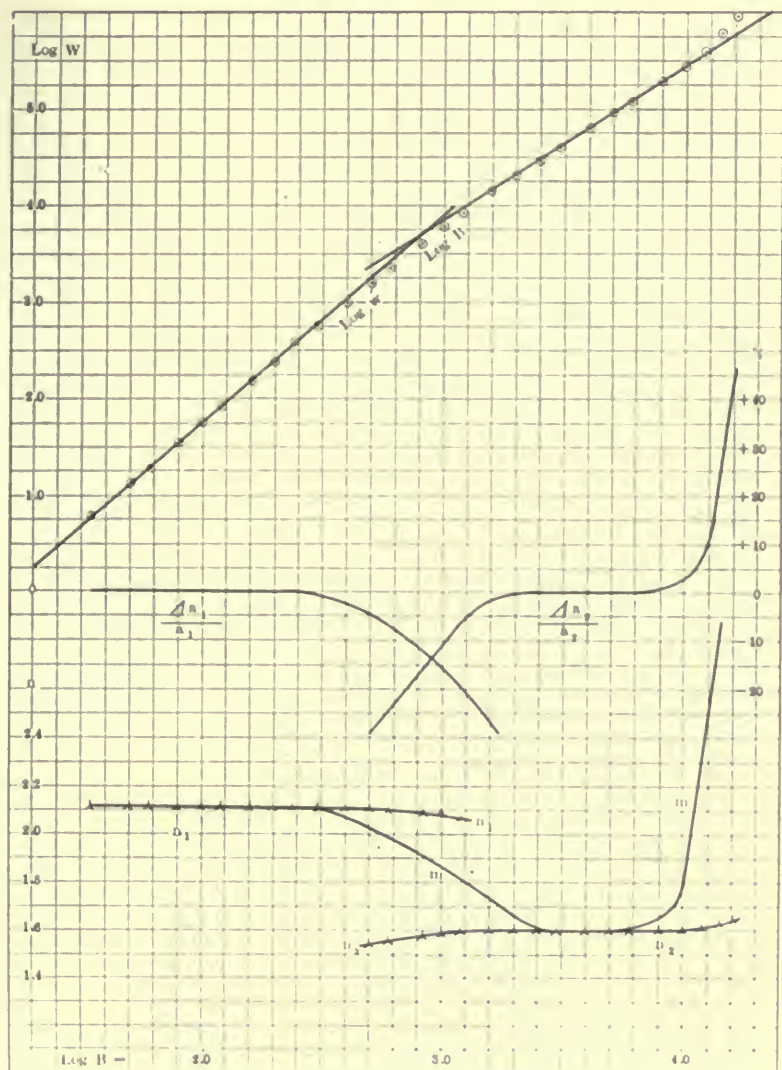


FIG. 84.

The percentage correction, which is to be applied to a_1 and a_2 respectively, to n_1 and n_2 , to make the formulas applicable

TABLE VIII.
HYSTERESIS OF SILICON STEEL.

B	w	$\log B$	$\log w$	$\frac{\Delta a_1}{a_1}$ %	$\frac{\Delta a_2}{a_2}$ %	$\frac{\Delta n_1}{n_1}$ %	$\frac{\Delta n_2}{n_2}$ %	n_1	n_2	m
35	6.4	1.544	0.806	+ 3.0	—	+ 0.46	—	2.120	—	—
50	13	1.699	1.117	— 0.2	—	— 0.03	—	2.109	—	2.03
60	19	1.778	1.279	— 1.1	—	— 0.13	—	2.107	—	2.14
80	36	1.903	1.556	+ 1.9	—	+ 0.20	—	2.114	—	2.12
100	57	2.000	1.752	— 0.2	—	— 0.02	—	2.110	—	2.09
120	83	2.079	1.922	+ 0.7	—	+ 0.07	—	2.111	—	2.16
160	156	2.204	2.193	+ 2.3	—	+ 0.21	—	2.114	—	2.10
200	245	2.301	2.389	+ 0.2	—	+ 0.02	—	2.110	—	2.07
250	394	2.398	2.595	+ 0.7	—	+ 0.06	—	2.111	—	2.03
300	571	2.477	2.757	— 0.5	—	— 0.04	—	2.109	—	1.98
400	1025	2.602	3.011	— 2.7	—	— 0.22	—	2.105	—	2.03
500	1610	2.699	3.207	— 4.7	— 29.5	— 0.37	— 3.52	2.102	1.544	2.02
600	2320	2.778	3.366	— 6.2	— 24.0	— 0.87	— 2.68	2.092	1.557	1.96
800	4030	2.903	3.605	— 16.5	— 15.8	— 1.17	— 1.72	2.085	1.573	1.91
1000	6150	3.000	3.789	— 15.7	— 11.1	— 1.14	— 1.06	2.086	1.583	1.89
1200	8680	3.079	3.938	— 18.7	— 6.0	— 2.02	— 0.55	2.087	1.5912	1.82
1600	14370	3.204	4.157	— 26.9	— 2.0	—	— 0.18	—	1.5971	1.73
2000	21000	3.301	4.322	—	0.0	—	0.00	—	1.6000	1.67
2500	30300	3.398	4.481	—	+ 0.9	—	+ 0.07	—	1.6011	1.62
3000	40500	3.477	4.607	—	+ 1.2	—	+ 0.09	—	1.6014	1.58
4000	63400	3.602	4.802	—	— 0.2	—	— 0.02	—	1.5997	1.58
5000	90600	3.699	4.957	—	— 0.2	—	— 0.02	—	1.5997	1.59
6000	120600	3.778	5.082	—	— 0.9	—	— 0.07	—	1.5989	1.61
8000	194100	3.903	5.288	—	+ 0.7	—	+ 0.05	—	1.6008	1.66
10000	282500	4.000	5.451	—	+ 2.6	—	+ 0.17	—	1.6027	1.77
12000	397500	4.079	5.599	—	+ 7.9	—	+ 0.50	—	1.6080	2.32
14000	609500	4.146	5.785	—	+ 27.7	—	+ 1.67	—	1.6267	2.88
16000	907500	4.204	5.958	—	+ 45.9	—	+ 2.85	—	1.6456	—
		$\Sigma_5 = 9.459$	7.626							
		$\Sigma_5 = 11.982$	12.945							
		$J = 2.523$	5.319							
		$n_1 = 2.523$	2.11							
		$\Sigma_{10} = 21.441$	20.571							
		$2.11 \times 21.441 = 45.241$								
		$J = -24.670$								
		$+ 10 = -2.467$								
		$= 7.533$	$= \log a_1$							
		$a_1 = 0.00341$								
		$\Sigma_4 = 13.380$	17.567							
		$\Sigma_4 = 14.982$	20.129							
		$J = 1.602$	2.562							
		$n_2 = 2.562$								
		$1.602 = 1.599 \simeq 1.60$								
		$\Sigma_8 = 28.362$	37.696							
		$1.60 \times 28.362 = 45.379$								
		$J = -7.683$								
		$+ 8 = -0.960$								
		$= 9.040 = \log a_2$								
		$a_2 = 0.1096$								

to the ranges where the logarithmic curve is not a straight line, are given in Table VIII as

$$\frac{\Delta a_1}{a_1}, \quad \frac{\Delta a_2}{a_2}, \quad \frac{\Delta n_1}{n_1}, \quad \frac{\Delta n_2}{n_2},$$

they are calculated as follows:

Assuming n as constant, $=n_0$, then a is not constant, $=a_0$, and the ratio:

$$\frac{\Delta a}{a} = \frac{a}{a_0} - 1$$

is the correction factor, and it is:

$$w = aB^{n_0},$$

hence:

$$\log w = \log a + n_0 \log B$$

and

$$\log a = \log w - n_0 \log B;$$

thus:

$$\log \frac{a}{a_0} = \log a - \log a_0 = \log w - \log a_0 - n_0 \log B,$$

and

$$\frac{\Delta a}{a} = \frac{a}{a_0} - 1 = N \log w - \log a_0 - n_0 \log B - 1. \quad (1)$$

Assuming a as constant, $=a_0$, then n is not constant, $=n_0$, and the ratio,

$$\frac{\Delta n}{n} = \frac{n}{n_0} - 1,$$

is the correction factor, and it is

$$w = a_0 B^n;$$

hence

$$\log w = \log a_0 + n \log B,$$

and

$$n \log B = \log w - \log a_0;$$

thus

$$\frac{n}{n_0} = \frac{n \log B}{n_0 \log B} = \frac{\log w - \log a_0}{n_0 \log B},$$

and

$$\frac{\Delta n}{n} = \frac{n}{n_0} - 1 = \frac{\log w - \log a_0 - n_0 \log B}{n_0 \log B} \quad . \quad . \quad . \quad (2)$$

by these equations (1) and (2) the correction factors in columns 5 to 8 of Table VIII are calculated, by using for a_0 and n_0 the values of the lower range curve, in columns 5 and 7, and the values of the medium range curve, in columns 6 and 8.

Thus, for instance, at $B = 1000$, the loss can be calculated by the equation,

$$w = a_1 B^{n_1},$$

by applying to a_1 the correction factor:

$$\begin{aligned} -15.7 \text{ per cent at constant: } n_1 = 2.11, \text{ that is,} \\ a_1 = 0.00341(1 - 0.157) = 0.00287; \end{aligned}$$

or by applying to n_1 the correction factor:

$$\begin{aligned} -1.14 \text{ per cent at constant: } a_1 = 0.00341, \text{ that is,} \\ n_1 = 2.11(1 - 0.0114) = 2.086. \end{aligned}$$

Or the loss can be calculated by the equation,

$$w = a_2 B^{n_2},$$

by applying to a_2 the correction factor:

$$\begin{aligned} -11.1 \text{ per cent at constant: } n_2 = 1.60, \text{ that is,} \\ a_2 = 0.1096(1 - 0.111) = 0.0974; \end{aligned}$$

or by applying to n_2 the correction factor:

$$\begin{aligned} -1.06 \text{ per cent at constant: } a_2 &= 0.1096, \text{ that is,} \\ n_2 &= 1.60(1 - 0.0106) = 1.583; \end{aligned}$$

and the loss may thus be given by either of the four expressions:

$$w = 0.00287B^{2.11} = 0.00341B^{2.086} = 0.0974B^{1.6} = 0.1096B^{1.587}.$$

As seen, the variation of the exponent n , required to extend the use of the parabolic equation into the range for which it does not strictly apply any more, is much less than the variation of the coefficient a , and a far greater accuracy is thus secured by considering the exponent n as constant—1.6 for medium and high values of B —and making the correction in coefficient a , outside of the range where the 1.6th power law holds rigidly.

In the last column of Table VIII is recorded the ratio of variation, $m = \frac{\int \log w}{\int \log B}$, as the averages each of two successive values. As seen, m agrees with the exponent n within the two ranges, where it is constant, but differs from it outside of these ranges. For instance, if B changes from 1600 downward, the ratio of variation m increases, while the exponent n slightly decreases.

In Fig. 84 are shown the percentage correction of the coefficients a_1 and a_2 , and also the two exponents n_1 and n_2 , together with the ratio of variation m .

The ratio of variation m is very useful in calculating the change of loss resulting from a small change of magnetic density, as the percentual change of loss w is m times the percentual (small) change of density.

As further example, the reader may reduce to empirical equations the series of observations given in Table IX. This table gives:

A. The candle-power L , as function of the power input p , of a 40-watt tungsten filament incandescent lamp.

B. The loss of power by corona (discharge into the air), p , in kw., in 1.895 km. of conductor, as function of the voltage e (in kv.) between conductor and return conductor, for the

TABLE IX.

A. Luminosity characteristic of 40-watt tungsten incandescent lamp.

 L = horizontal candlepower. p = watts input.

L	p	L	p	L	p	L	p	L	p
2	12.25	20	31.64	40	44.14	128	76.77	382	135.6
4	16.33	24	34.55	44	45.42	192	95.24	460	145.2
8	21.35	28	37.29	48	47.05	256	109.0	—	—
12	25.60	32	39.26	64	54.31	291	118.2	—	—
16	28.91	36	41.47	96	65.73	320	123.1	—	—

B. Corona loss of high-voltage transmission line; at 60 cycles:

1895 m., length of conductor.

3.10 m., distance between conductors.

No. 000 seven-strand cable, 1.18 cm. diameter.

-13° C.; 76.2 cm. barometer; sunshine.

 e = kilovolts between conductors, effective. p = kilowatts loss.

e	p	e	p	e	p	e	p	e	p
79.8	0.01	141.5	0.09	181.0	1.02	221.0	8.70	212.0	6.44
90.7	0.01	147.0	0.08	186.2	1.55	227.0	10.66	219.0	8.31
101.5	0.02	153.6	0.12	192.6	2.49	234.0	13.25	—	—
109.5	0.03	159.0	0.16	200.6	3.77	189.0	2.10	—	—
120.5	0.04	169.8	0.35	208.6	5.34	195.0	2.88	—	—
130.0	0.06	174.0	0.53	216.0	7.13	203.8	4.72	—	—

C. Volume-pressure characteristic of dry steam at its boiling-point.

 t = degrees C. P = pressure, in kg. per cm.² V = volume, in m.³ per kg.

t	P	V	t	P	V	t	P	V
59.8	0.2	7.806	132.8	3.0	0.612	169.5	8.0	0.244
80.9	0.5	3.297	142.8	4.0	0.467	178.9	10.0	0.197
99.1	1.0	1.717	151.0	5.0	0.379	186.9	12.0	0.167
119.6	2.0	0.896	157.9	6.0	0.319	197.2	15.0	0.135

distance of 310 cm. between the conductors, and the conductor diameter of 1.18 cm.

C. The relation between steam pressure P , in kg. per cm.², and the steam volume V , in m.³, at the boiling-point, per kg. of dry steam.

D. Periodic Curves.

158. All periodic functions of time or distance can be expressed by a trigonometric series, or Fourier series, as has been discussed in Chapter III, and the methods of resolution, and the arrangements to carry out the work rapidly, have also been discussed in Chapter III.

The resolution of a periodic function thus consists in the determination of the higher harmonics, which are superimposed on the fundamental wave.

As periodic functions are of the greatest importance in electrical engineering, in the theory of alternating current phenomena, a familiarity with the wave shapes produced by the different harmonics is desirable. This familiarity should be sufficient to enable one in most cases to judge immediately from the shape of the wave, as given by oscillograph, etc., on the harmonics which are present or at least which predominate.

The effect of the lower harmonics, such as the third, fifth, etc., (or the second, fourth, etc., where present), is to change the shape of the wave, make it differ from sine shape, giving such features as flat top wave, peaked wave, saw-tooth, double and triple peaked, steep zero, flat zero, etc., while the high harmonics do not change the shape of the wave so much, as superimpose ripples on it.

ODD LOWER HARMONICS.

159. To elucidate the variation in shape of the alternating waves caused by various lower harmonics, superimposed upon the fundamental at different relative positions, that is, different phase angles, in Figs. 85 and 86 are shown the effect of a third harmonic, of 10 per cent and 30 per cent of the fundamental, respectively. *A* gives the fundamental, and *C D E F G* the waves resulting by the superposition of the triple harmonic in phase with the fundamental (*C*), under 45 deg. lead (*D*), 90 deg. lead or quadrature (*E*), 135 deg. lead (*F*) and opposition

(G). (The phase differences here are referred to the maximum of the fundamental: with waves of different frequencies, the phase differences naturally change from point to point, and in speaking of phase difference, the reference point on the wave

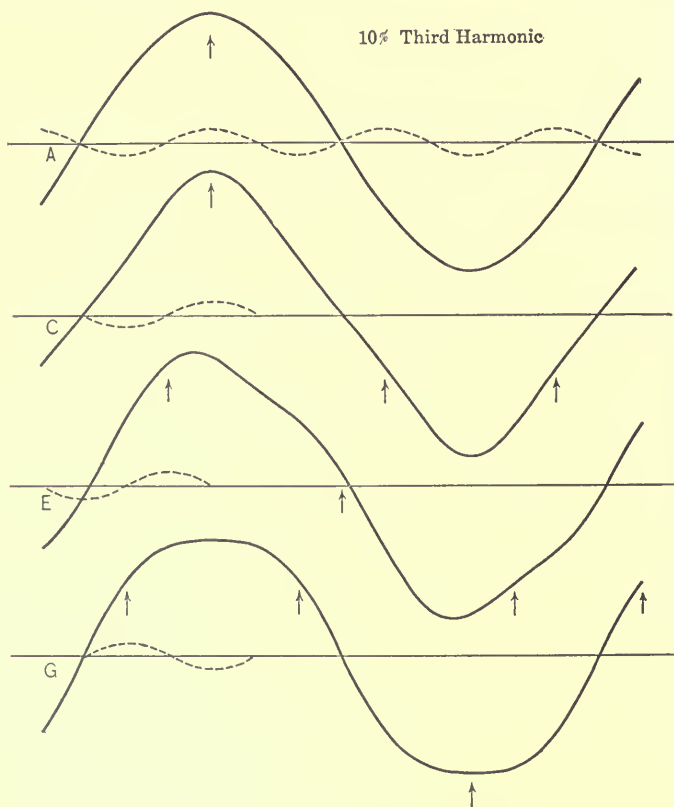


FIG. 85. Effect of Small Third Harmonic.

must thus be given. For instance, in *C* the third harmonic is in phase with the fundamental at the maximum point of the latter, but in opposition at its zero point.)

The equations of these waves are:

$$A: y = 100 \cos \beta$$

$$C: y = 100 \cos \beta + 10 \cos 3\beta$$

$$E: y = 100 \cos \beta + 10 \cos (3\beta + 90 \text{ deg.})$$

$$G: y = 100 \cos \beta + 10 \cos (3\beta + 180 \text{ deg.})$$

$$= 100 \cos \beta - 10 \cos 3\beta$$

$$C: y = 100 \cos \beta + 30 \cos 3\beta$$

$$D: y = 100 \cos \beta + 30 \cos (3\beta + 45 \text{ deg.})$$

$$E: y = 100 \cos \beta + 30 \cos (3\beta + 90 \text{ deg.})$$

$$F: y = 100 \cos \beta + 30 \cos (3\beta + 135 \text{ deg.})$$

$$G: y = 100 \cos \beta + 30 \cos (3\beta + 180 \text{ deg.})$$

$$= 100 \cos \beta - 30 \cos 3\beta$$

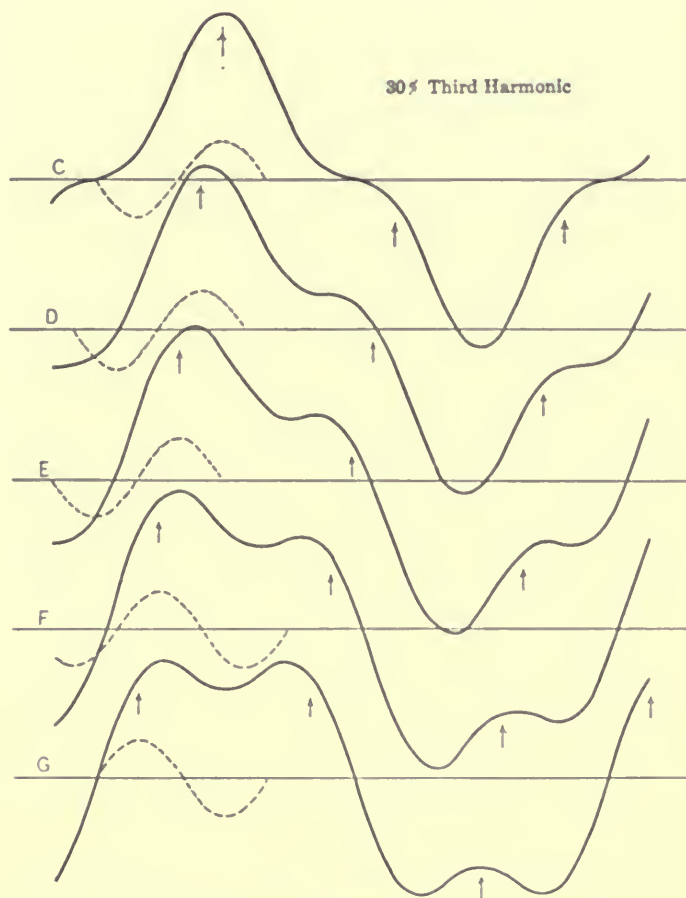


FIG. 86. Effect of Large Third Harmonic.

In all these waves, one cycle of the triple harmonic is given in dotted lines, to indicate its relative position and intensity, and the maxima of the harmonics are indicated by the arrows.

As seen, with the harmonic in phase or in opposition (C and G), the waves are symmetrical; with the harmonic out of phase, the waves are unsymmetrical, of the so-called "saw tooth" type, and the saw tooth is on the rising side of the wave with a lagging, on the decreasing side with a leading triple harmonic.

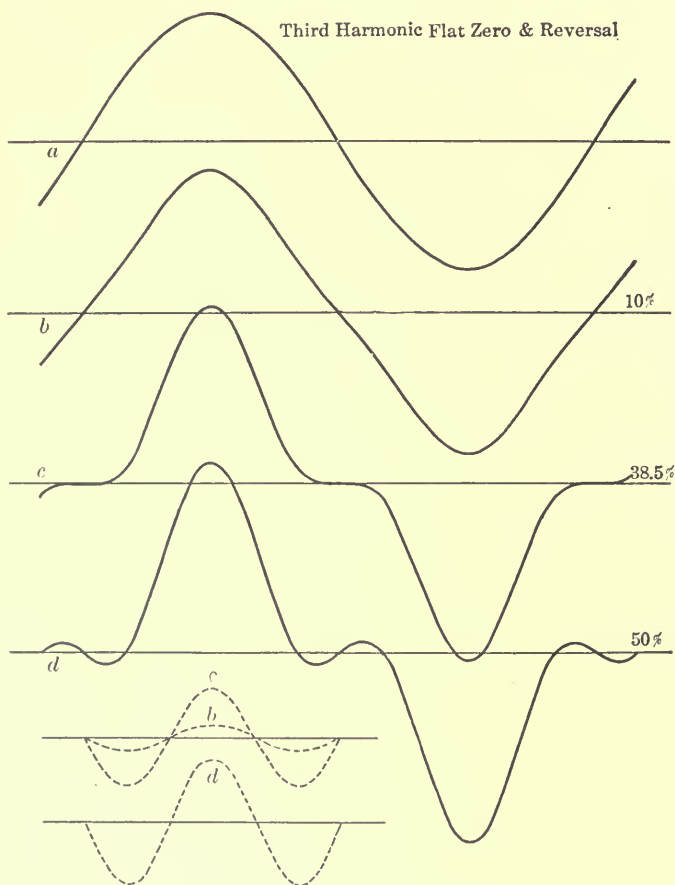


FIG. 87. Flat Zero and Reversal by Third Harmonic.

The latter are shown in D , E , F ; the former have the same shape but reversed, that is, rising and decreasing side of the wave interchanged, and therefore are not shown.

The triple harmonic in phase with the fundamental, C , gives a peaked wave with flat zero, and the peak and the flat zero

become the more pronounced, the higher the third harmonic, until finally the flat zero becomes a double reversal of voltage, as shown in Fig. 87*d*.

Fig. 87 shows the effect of a gradual increase of an in-phase triple harmonic: *a* is the fundamental, *b* contains a 10 per

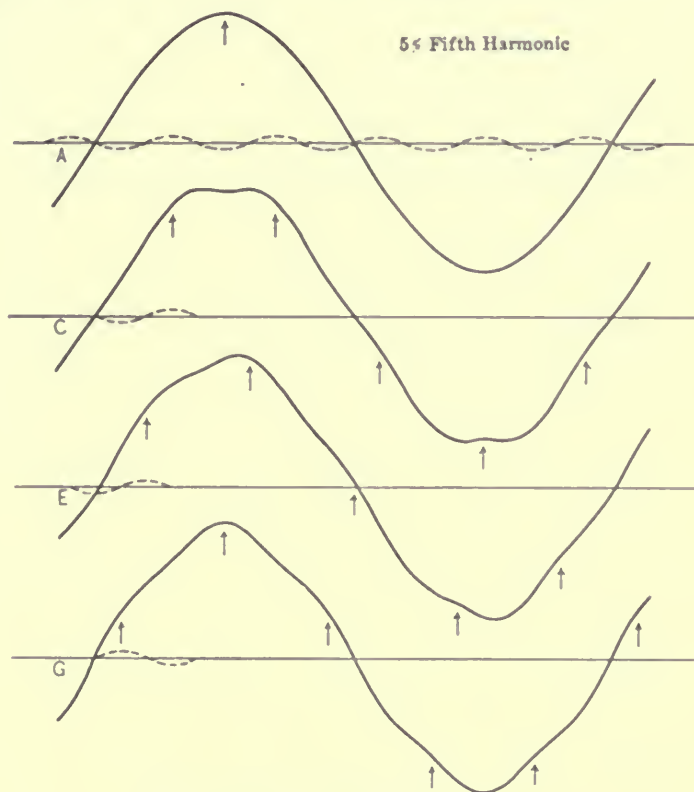


FIG. 88. Effect of Small Fifth Harmonic.

cent, *c* a 38.5 per cent and *d* a 50 per cent triple harmonic, as given by the equations:

$$a: y = 100 \cos \beta$$

$$b: y = 100 \cos \beta + 10 \cos 3\beta$$

$$c: y = 100 \cos \beta + 38.5 \cos 3\beta$$

$$d: y = 100 \cos \beta + 50 \cos 3\beta$$

At *c*, the wave is entirely horizontal at the zero, that is, remains zero for an appreciable time at the reversal. In this figure, the three harmonics are shown separately in dotted lines, in their relative intensities.

A triple harmonic in opposition to the fundamental (Figs. 85 and 86*G*) is characterized by a flat top and steep zero, and

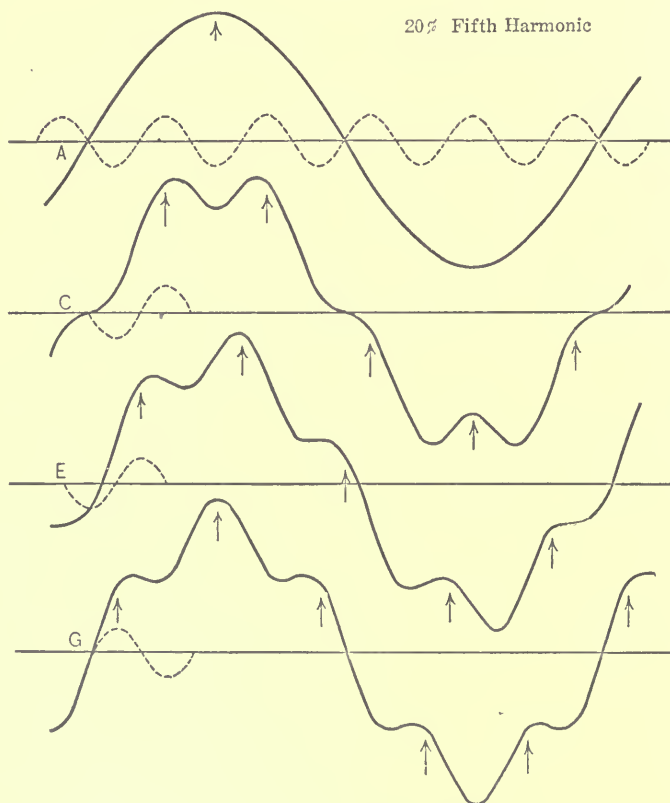


FIG. 89. Effect of Large Fifth Harmonic.

with the increase of the third harmonic, the flat top develops into a double peak (Fig. 86*G*), while steepness at the point of reversal increases.

The simple saw tooth, produced by a triple harmonic in quadrature with the fundamental is shown in Fig. 85*E*. With increasing triple harmonic, the hump of the saw tooth becomes

more pronounced and changes to a second and lower peak, as shown in Fig. 86. This figure gives the variation of the saw-tooth shape from 45 to 45 deg. phase difference: With the phase of the third harmonic shifting from in-phase to 45 deg. lead, the flat zero, by moving up on the wave, has formed a hump or saw tooth low down on the decreasing (and with 45 deg. lag on the increasing) side of the wave. At 90 deg. lead, the saw tooth has moved up to the middle of the down branch of the wave, and with 135 deg. lead, has moved still further up, forming practically a second, lower peak. With 180 deg. lead—or opposition of phase—the hump of the saw tooth has moved up to the top, and formed the second peak—or the flat top, with a lower third harmonic, as in Fig. 85*G*.

Figs. 88 and 89 give the effect of the fifth harmonic, superimposed on the fundamental, of 5 per cent in Fig. 88, and of 20 per cent in Fig. 89. Again *A* gives the fundamental sine wave, *C* the effect of the fifth harmonic in opposition with the fundamental, *E* in quadrature (lagging) and *G* in phase. One cycle of the fifth harmonic is shown in dotted lines, and the maxima of the harmonics indicated by the arrows.

The equations of these waves are given by:

$$\begin{aligned} A: y &= 100 \cos \beta \\ C: y &= 100 \cos \beta - 5 \cos 5\beta \\ E: y &= 100 \cos \beta - 5 \cos (5\beta + 90 \text{ deg.}) \\ G: y &= 100 \cos \beta + 5 \cos 5\beta \end{aligned}$$

$$\begin{aligned} A: y &= 100 \cos \beta \\ C: y &= 100 \cos \beta - 20 \cos 5\beta \\ E: y &= 100 \cos \beta - 20 \cos (5\beta + 90 \text{ deg.}) \\ G: y &= 100 \cos \beta + 20 \cos 5\beta \end{aligned}$$

In the distortion caused by the fifth harmonic (in opposition to the fundamental) flat top (Fig. 88*C*) or double peak (at higher values of the harmonic, Fig. 89*C*), is accompanied by flat zero (or, at very high values of the fifth harmonic, double reversal at the zero, similar as in Fig. 87*d*), while in the distortion by the third harmonic it is accompanied by sharp zero.

With the fifth harmonic in phase with the fundamental, a peaked wave results with steep zero, Fig. 88*G*, and the transi-

tion from the steep zero to the peak, with larger values of the fifth harmonic, then develops into two additional peaks, thus giving a treble peaked wave, Fig. 88G, with steep zero. The beginning of treble peakedness is noticeable already in Fig. 88G, with only 5 per cent of fifth harmonic.

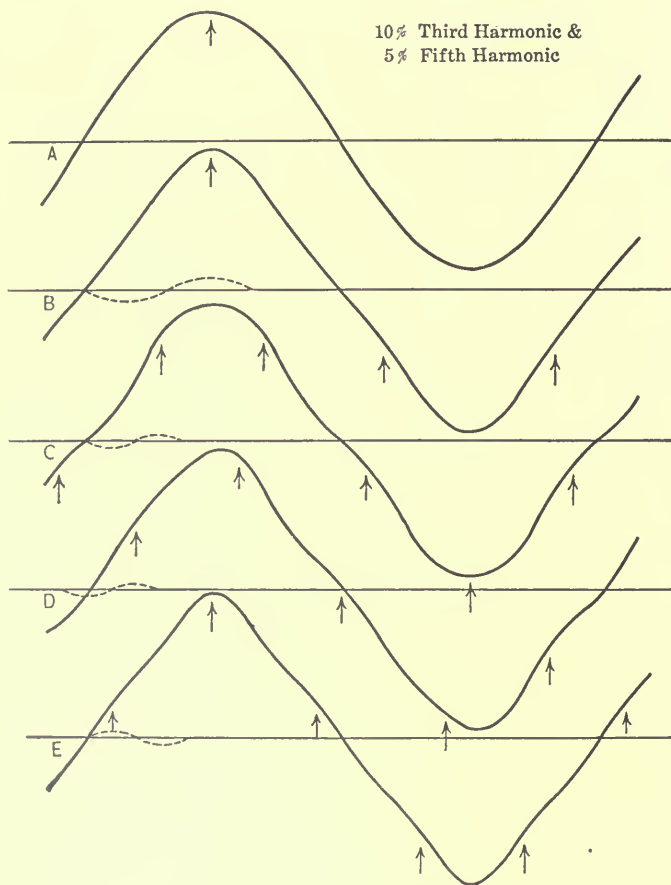


FIG. 90. Third and Fifth Harmonic.

With the seventh harmonic, the treble-peaked wave would be accompanied by flat zero, and a quadruple-peaked wave would give steep zero (Fig. 95).

The fifth harmonic out of phase with the fundamental again gives saw-tooth waves, Figs. 88 and 89E, but the saw tooth

produced by the fifth harmonic contains two humps, that is, is double, with one hump low down, and the other high up on the curve, thereby giving the transition from the symmetrical double peak *C* to the symmetrical treble peak *G*.

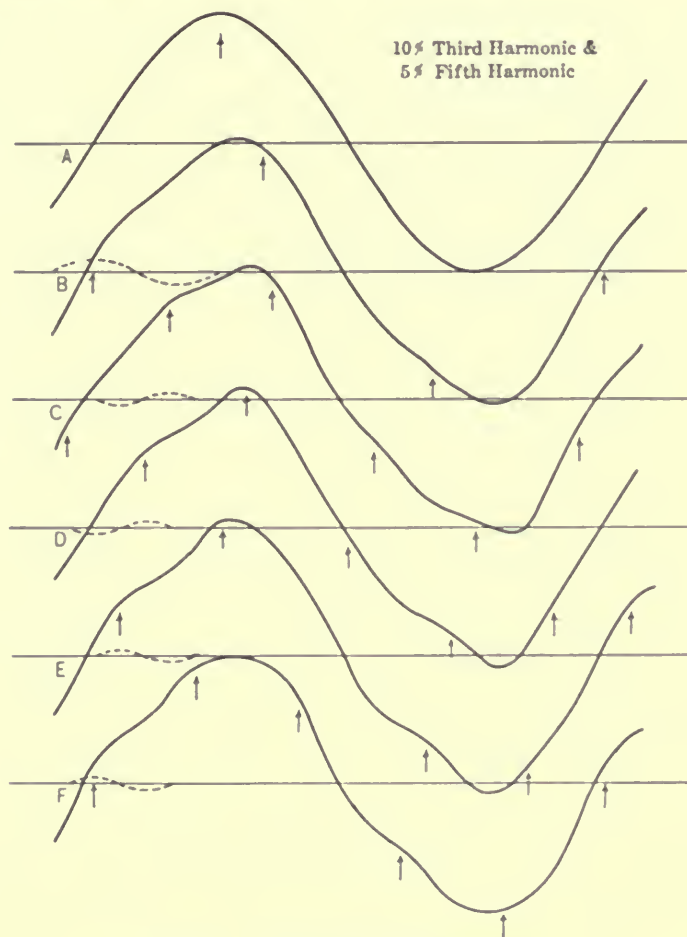


FIG. 91. Third and Fifth Harmonic.

160. Characteristic of the effect of the third harmonic thus is:

Coincidence of peak with flat zero or double reversal, of steep zero with flat top or double peak, and single hump or saw tooth.

While characteristic of the effect of the fifth harmonic is:

Coincidence of peak with steep zero, or treble peak, of flat top or double peak with flat zero or double reversal, and double saw-tooth.

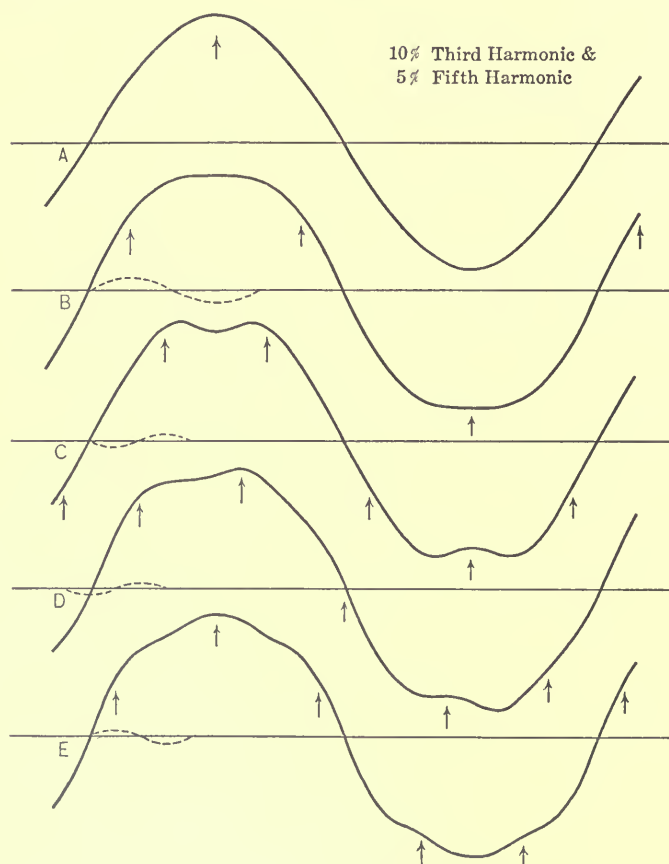


FIG. 92. Third and Fifth Harmonic.

By thus combining third and fifth harmonics of proper values, they can be made to neutralize each other's effect in any one of their characteristics, but then accentuate each other in the other characteristic.

Thus peak and flat zero of the triple harmonic combined with peak and steep zero of the fifth harmonic, gives a peaked wave with normal sinusoidal appearance at the zero value; combin-

ing the flat tops or double peaks of both harmonics, the flat zero of the one neutralizes the steep zero of the other, and we get a flat top or double peak with normal zero. Or by combining the peak of the third harmonic with the flat top of the fifth we get a wave with normal top, but steep zero, and we get a wave with normal top, but flat zero or double reversal, by combining the triple harmonic peak with the fifth harmonic flat top.

Thus any of the characteristics can be produced separately by the combination of the third and fifth harmonic.

By combining third and fifth harmonics out of phase with fundamental—such as give single or double saw-tooth shapes, the various other saw-tooth shapes are produced, and still further saw-tooth shapes, by combining a symmetrical (in phase or in opposition) third harmonic with an out of phase fifth, or inversely.

These shapes produced by the superposition, under different phase angles, of fifth and third harmonics on the fundamental, and their gradual change into each other by the shifting in phase of one of the harmonics, are shown in Figs. 90, 91 and 92 for a third harmonic of 10 per cent, and a fifth harmonic of 5 per cent of the fundamental.

In Fig. 90 the third harmonic is in phase, in Fig. 91 in quadrature lagging, and in Fig. 92 in opposition with the fundamental. *A* gives the fundamental, *B* the fundamental with the third harmonic only, and *C*, *D*, *E*, *F*, the waves resulting from the superposition of the fifth harmonic on the combination of fundamental and third harmonic, given as *B*. In *C* the fifth harmonic is in opposition, in *D* in quadrature lagging, in *E* in phase, and in *F* in quadrature leading.

We see here round tops with flat zero (Fig. 90*C*), nearly triangular waves (Fig. 90*E*), approximate half circles (Fig. 92*E*), sine waves with a dent at the top (Fig. 92*C*), and various different forms of saw tooth.

The equations of these waves are:

$$A: y = 100 \cos \beta$$

$$B: y = 100 \cos \beta + 10 \cos 3\beta$$

$$C: y = 100 \cos \beta + 10 \cos 3\beta - 5 \cos 5\beta$$

$$D: y = 100 \cos \beta + 10 \cos 3\beta - 5 \cos (\beta \mp 90 \text{ deg.})$$

$$E: y = 100 \cos \beta + 10 \cos 3\beta + 5 \cos 5\beta$$

$$A: y = 100 \cos \beta$$

$$B: y = 100 \cos \beta - 10 \cos (3\beta + 90 \text{ deg.})$$

$$C: y = 100 \cos \beta - 10 \cos (3\beta + 90 \text{ deg.}) - 5 \cos 5\beta$$

$$D: y = 100 \cos \beta - 10 \cos (3\beta + 90 \text{ deg.}) - 5 \cos (5\beta + 90 \text{ deg.})$$

$$E: y = 100 \cos \beta - 10 \cos (3\beta + 90 \text{ deg.}) + 5 \cos 5\beta$$

$$F: y = 100 \cos \beta - 10 \cos (3\beta + 90 \text{ deg.}) + 5 \cos (5\beta + 90 \text{ deg.})$$

$$A: y = 100 \cos \beta$$

$$B: y = 100 \cos \beta - 10 \cos 3\beta$$

$$C: y = 100 \cos \beta - 10 \cos 3\beta - 5 \cos 5\beta$$

$$D: y = 100 \cos \beta - 10 \cos 3\beta - 5 \cos (5\beta + 90 \text{ deg.})$$

$$E: y = 100 \cos \beta - 10 \cos 3\beta + 5 \cos 5\beta$$

EVEN HARMONICS.

161. Characteristic of the wave-shape distortion of even harmonics is that the wave is not a symmetrical wave, but the two half waves have different shapes, and the characteristics of the negative half wave are opposite to those of the positive. This is to be expected, as an even harmonic, which is in phase with the positive half wave of the fundamental, is in opposition with the negative; when leading in the positive, it is lagging in the negative, and inversely.

Fig. 93 shows the effect of a second harmonic of 30 per cent of the fundamental *A*, superimposed in quadrature, 60 deg. phase displacement, 30 deg. displacement and in phase, in *B*, *C*, *D* and *E* respectively.

The equations of these waves are:

$$A: y = 100 \cos \beta \text{ and } y' = 30 \cos (2\beta - 90)$$

$$B: y = 100 \cos \beta + 30 \cos (2\beta - 90)$$

$$C: y = 100 \cos \beta + 30 \cos (2\beta - 60)$$

$$D: y = 100 \cos \beta + 30 \cos (2\beta - 30)$$

$$E: y = 100 \cos \beta + 30 \cos 2\beta$$

Quadrature combination (Fig. 93*B*) gives a wave where the rising side is flat, the decreasing side steep, and inversely with the other half wave. *C* and *D* give a peaked wave for the one, a saw tooth for the other half wave, and *E*, coincidence of phase of fundamental and second harmonic, gives a combination of one peaked half wave with one flat-top or double-peaked wave.

Characteristic of *C*, *D* and *E* is, that the two half waves are of unequal length.

In general, even harmonics, if of appreciable value are easily recognized by the difference in shape, of the two half waves.

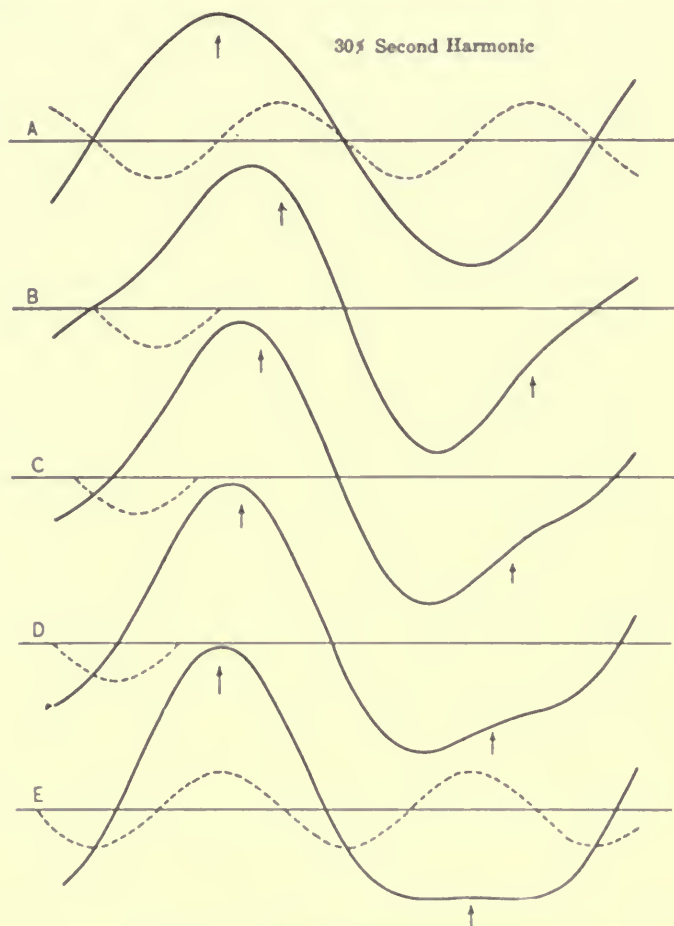


FIG. 93. Effect of Second Harmonic.

By the combination of the second harmonic with the third harmonic (or the fifth), some of the features can be intensified, others suppressed.

An illustration hereof is shown in Fig. 94 in the production

of a wave, in which the one half wave is a short high peak, the other a long flat top, by the superposition of a second harmonic of 46.5 per cent, and a third harmonic of 10 per cent both in phase with the fundamental.

A gives the fundamental sine wave, *B* and *C* the second and third harmonic, *D* the combination of fundamental and second

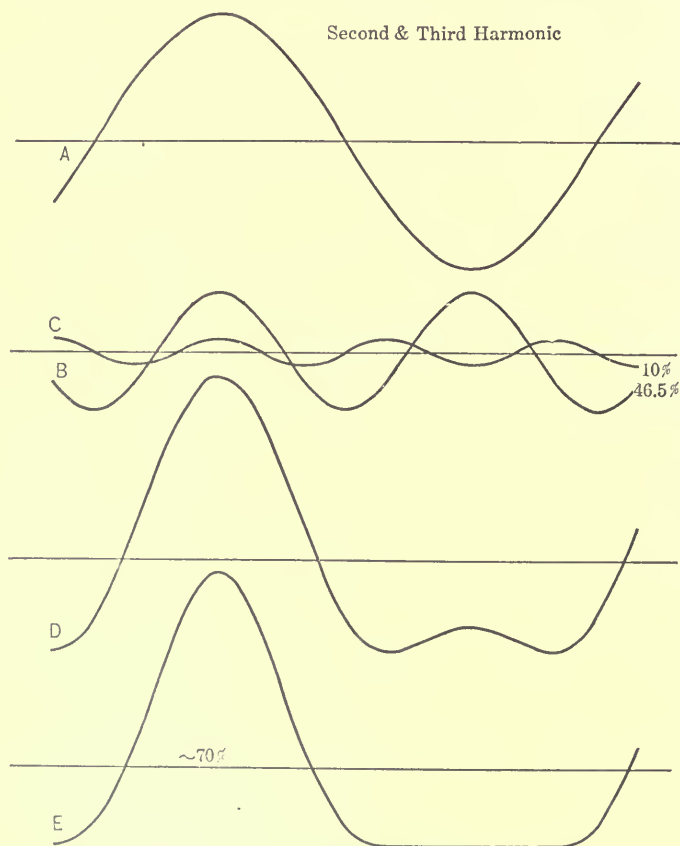


FIG. 94. Peak and Flat Top by Second and Third Harmonic.

harmonic, giving a double peaked negative half wave, and *E* the addition of the third harmonic to the wave *D*. Thereby the double peak of the negative half wave is flattened to a long flat top, and the peak of the positive half wave intensified and shortened, so that the positive maximum is about two and one-

half times the negative maximum, and the negative half wave nearly 75 per cent longer than the positive half wave.

The equations of these waves are given by:

$$A: y = 100 \cos \beta$$

$$B: y = 46.5 \cos 2\beta$$

$$C: y = 10 \cos 3\beta$$

$$D: y = 100 \cos \beta + 46.5 \cos 2\beta$$

$$E: y = 100 \cos \beta + 46.5 \cos 2\beta + 10 \cos 3\beta$$

HIGH HARMONICS.

162. Comparing the effect of the fifth harmonic, Figs. 88 and 89, with that of the third harmonic, Figs. 85 and 86, it is seen

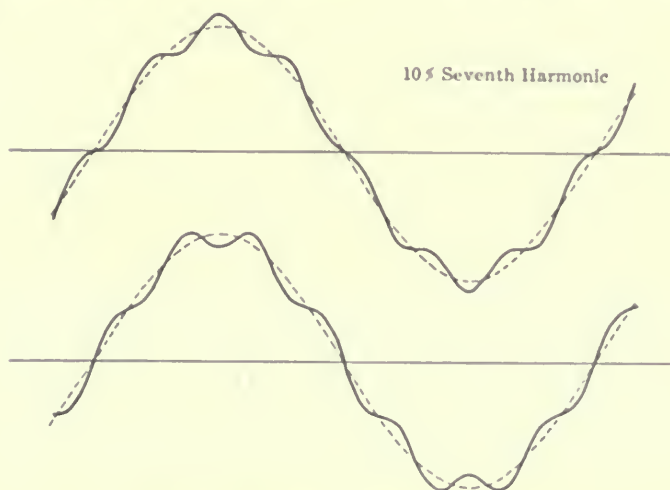


FIG. 95. Effect of Seventh Harmonic.

that a fifth harmonic, even if very small, is far easier distinguished, that is, merges less into the fundamental than the third harmonic. Still more this is the case with the seventh harmonic, as shown in Fig. 95 in phase and in opposition, of 10 per cent intensity. This is to be expected: sine waves which do not differ very much in frequency, such as the fundamental and the second or third harmonic, merge into each other and form a resultant shape, a distorted wave of characteristic appearance,

while sine waves of very different frequencies, as the fundamental and its eleventh harmonic, in Fig. 96, when superimposed, remain distinct from each other; the general shape of the wave is the fundamental sine, and the high harmonics appear as ripples upon the fundamental, thus giving what may be called a corrugated sine wave. By counting the number of ripples per

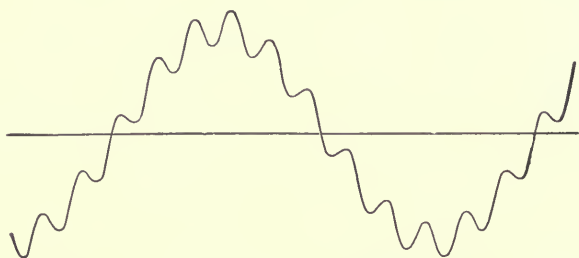


FIG. 96. Wave in which Eleventh Harmonic Predominates.

complete wave, or per half wave, the order of the harmonic can then rapidly be determined. For instance, the wave shown in Fig. 96 contains mainly the eleventh harmonic, as there are eleven ripples per wave. The wave shown by the oscillogram Fig. 97 shows the twenty-third harmonic, etc.



FIG. 97. C D 23510. Alternator Wave with Single High Harmonic.

Very frequently high harmonics appear in pairs of nearly the same frequency and intensity, as an eleventh and a thirteenth harmonic, etc. In this case, the ripples in the wave shape show maxima, where the two harmonics coincide, and nodes, where the two harmonics are in opposition. The presence of nodes makes the counting of the number of ripples per complete wave more difficult. A convenient method of procedure in this case

is, to measure the distance or space between the maxima of one or a few ripples in the range where they are pronounced, and count the number of nodes per cycle. For instance, in the wave, Fig. 98, the space of two ripples is about 60 deg., and two nodes exist per complete wave. 60 deg. for two ripples, give

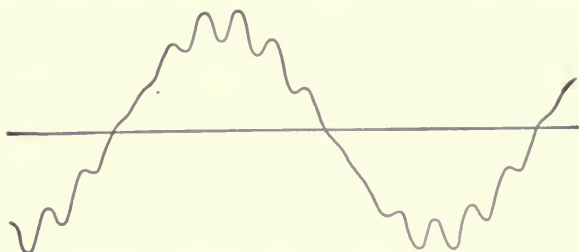


FIG. 98. Wave in which Eleventh and Thirteenth Harmonics Predominate.

$2 \times \frac{360}{60} = 12$ ripples per complete wave, as the average frequency of the two existing harmonics, and since these harmonics differ by 2 (since there are two nodes), their order is the eleventh and the thirteenth harmonics.

163. This method of determining two similar harmonics, with

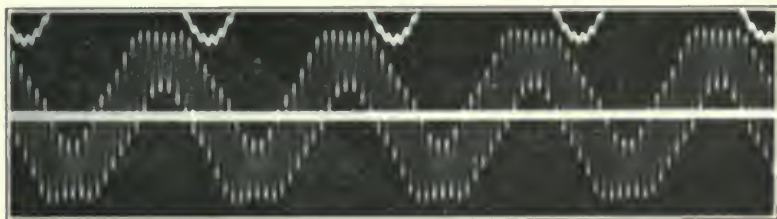


FIG. 99. CD 23512. Alternator Wave with Two Very Unequal High Harmonics.

a little practice, becomes very convenient and useful, and may frequently be used visually also, in determining the frequency of hunting of synchronous machines, etc. In the phenomenon of hunting, frequently two periods are superimposed, a forced frequency, resulting from speed of generator, etc., and the natural frequency of the machine. Counting the number of impulses, f , per minute, and the number of nodes, n , gives the

two frequencies: $f + \frac{n}{2}$ and $f - \frac{n}{2}$; and as one of these frequencies is the impressed engine frequency, this affords a check.

Where the two high harmonics of nearly equal order, as the eleventh and the thirteenth in Fig. 98, are approximately equal in intensity, at the nodes the ripples practically disappear, and between the nodes the ripples give a frequency intermediate between the two components: Apparently the twelfth harmonic in Fig. 98. In this case the two constituents are easily determined: $12 - 1 = 11$, and $12 + 1 = 13$.

Where of the two constituents one is greater than the other the wave still shows nodes, but the ripples do not entirely disappear at the nodes, but merely decrease, that is, the wave shows a sine with ripples which increase and decrease along the waves

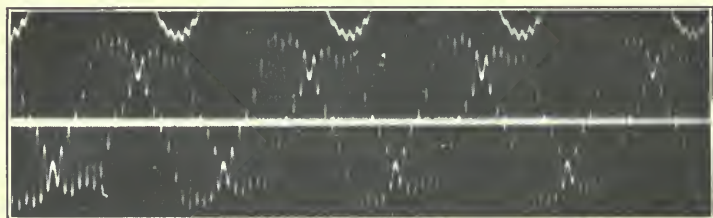


FIG. 100. CD 23511. Alternator Wave with Two Nearly Equal High Harmonics.

as shown by the oscillograms 99 and 100. In this case, one of the two high frequencies is given by counting the total number of ripples, but it may at first be in doubt, whether the other component is higher or lower by the number of nodes. The decision then is made by considering the length of the ripple at the node: If the length is a maximum at the node, the secondary harmonic is of higher frequency than the predominating one; if the length of the ripple at the node is a minimum, the secondary frequency is lower than the predominating one. This is illustrated in Fig. 101. In this figure, *A* and *B* represent the tenth and twelfth harmonic of a wave, respectively; *C* gives their superposition with the lower harmonic *A* predominating, while *B* is only of half the intensity of *A*. *D* gives the superposition of *A* and *B* at equal intensity, and *E* gives the superposition with the higher frequency *B* predominating. That is, the respective equations would be:

$$A: y = \cos 10\beta$$

$$B: y = \cos 12\beta$$

$$C: y = \cos 10\beta + 0.5 \cos 12\beta$$

$$D: y = \cos 10\beta + \cos 12\beta$$

$$E: y = 0.5 \cos 10\beta + \cos 12\beta$$

As seen, in *C* the half wave at the node is abnormally long, showing the preponderance of the lower frequency, in *E* abnor-

Superposition of High Harmonics

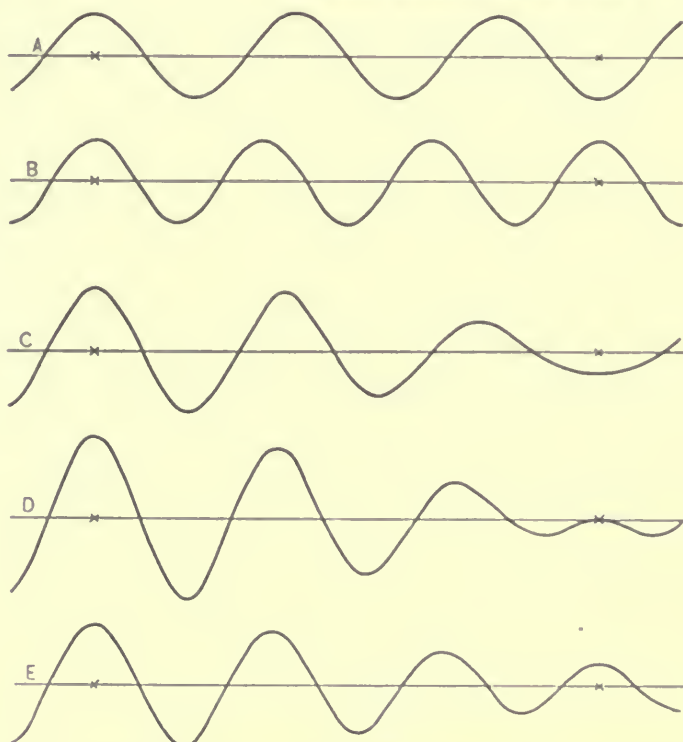


FIG. 101. Superposition of Two High Harmonics of Various Intensities.

mally short, showing the preponderance of the higher frequency.

In alternating-current and voltage waves, the appearance of two successive high harmonics is quite frequent. For instance, if an alternating current generator contains n slots per pole, this produces in the voltage wave the two harmonics of orders

$2n-1$ and $2n+1$. Such is the origin of the harmonics in the oscillograms Figs. 99 and 100.

The nature of the increase and decrease of the ripples and the formation of the nodes by the superposition of two adjacent high harmonics is best seen by combining their expressions trigonometrically.

Thus the harmonics:

$$y_1 = \cos (2n-1)\beta$$

and $y_2 = \cos (2n+1)\beta$

combined give the resultant:

$$\begin{aligned} y &= y_1 + y_2 \\ &= \cos (2n-1)\beta + \cos (2n+1)\beta \\ &= 2 \cos \beta \cos 2n\beta \end{aligned}$$

that is, give a wave of frequency $2n$ times the fundamental: $\cos 2n\beta$, but which is not constant, but varies in intensity by the factor $2 \cos \beta$.

Not infrequently wave-shape distortions are met, which are not due to higher harmonics of the fundamental wave, but are incommensurable therewith. In this case there are two entirely unrelated frequencies. This, for instance, occurs in the secondary circuit of the single-phase induction motor; two sets of currents, of the frequencies f_s and $(2f-f_s)$ exist (where f is the primary frequency and f_s the frequency of slip). Of this nature, frequently, is the distortion produced by surges, oscillations, arcing grounds, etc., in electric circuits; it is a combination of the natural frequency of the circuit with the impressed frequency. Telephonic currents commonly show such multiple frequencies, which are not harmonics of each other.

CHAPTER VII.

NUMERICAL CALCULATIONS.

164. Engineering work leads to more or less extensive numerical calculations, when applying the general theoretical investigation to the specific cases which are under consideration. Of importance in such engineering calculations are:

- (a) The method of calculation.
- (b) The degree of exactness required in the calculation.
- (c) The intelligibility of the results.
- (d) The reliability of the calculation.

a. Method of Calculation.

Before beginning a more extensive calculation, it is desirable carefully to scrutinize and to investigate the method, to find the simplest way, since frequently by a suitable method and system of calculation the work can be reduced to a small fraction of what it would otherwise be, and what appear to be hopelessly complex calculations may thus be carried out quickly and expeditiously by a proper arrangement of the work. Indeed, the most important part of engineering work—and also of other scientific work—is the determination of the method of attacking the problem, whatever it may be, whether an experimental investigation, or a theoretical calculation. It is very rarely that important problems can be solved by a direct attack, by brutally forcing a solution, and then only by wasting a large amount of work unnecessarily. It is by the choice of a suitable method of attack, that intricate problems are reduced to simple phenomena, and then easily solved; frequently in such cases requiring no solution at all, but being obvious when looked at from the proper viewpoint.

Before attacking a more complicated problem experimentally or theoretically, considerable time and study should thus first be devoted to the determination of a suitable method of attack.

The next then, in cases where considerable numerical calculations are required, is the method of calculation. The most convenient one usually is the arrangement in tabular form.

As example, consider the problem of calculating the regulation of a 60,000-volt transmission line, of $r=60$ ohms resistance, $x=135$ ohms inductive reactance, and $b=0.0012$ condensive susceptance, for various values of non-inductive, inductive, and condensive load.

Starting with the complete equations of the long-distance transmission line, as given in "Theory and Calculation of Transient Electric Phenomena and Oscillations," Section III, paragraph 9, and considering that for every one of the various power-factors, lag, and lead, a sufficient number of values have to be calculated to give a curve, the amount of work appears hopelessly large.

However, without loss of engineering exactness, the equation of the transmission line can be simplified by approximation, as discussed in Chapter V, paragraph 123, to the form,

$$\left. \begin{aligned} E_1 &= E_0 \left\{ 1 + \frac{ZY}{2} \right\} + ZI_0 \left\{ 1 + \frac{ZY}{6} \right\}; \\ I_1 &= I_0 \left\{ 1 + \frac{ZY}{2} \right\} + YE_0 \left\{ 1 + \frac{ZY}{6} \right\}, \end{aligned} \right\} \quad \dots \quad (1)$$

where E_0 , I_0 are voltage and current, respectively at the step-down end, E_1 , I_1 at the step-up end of the line; and

$Z=r+jx=60+135j$ is the total line impedance;

$Y=g+jb=+0.0012j$ is the total shunted line admittance.

Herefrom follow the numerical values:

$$\begin{aligned} 1 + \frac{ZY}{2} &= 1 + \frac{(60+135j)(+0.0012j)}{2} \\ &= 1 + 0.036j - 0.081 = 0.919 + 0.036j; \end{aligned}$$

$$1 + \frac{ZY}{6} = 1 + 0.012j - 0.027 = 0.973 + 0.012j;$$

$$Z \left\{ 1 + \frac{ZY}{6} \right\} = (60 + 135j)(0.973 + 0.012j) \\ = 58.4 + 0.72j + 131.1j - 1.62 = 56.8 + 131.8j;$$

$$Y \left\{ 1 + \frac{ZY}{6} \right\} = (+0.0012j)(0.973 + 0.012j) \\ = +0.001168j - 0.0000144 = (-0.0144 + 1.168j)10^{-3}$$

hence, substituting in (1), the following equations may be written:

$$\left. \begin{aligned} E_1 &= (0.919 + 0.036j)E_0 + (56.8 + 131.8j)I_0 = A + B \\ I_1 &= (0.919 + 0.036j)I_0 - (0.0144 - 1.168j)E_0 10^{-3} = C - D. \end{aligned} \right\} (2)$$

165. Now the work of calculating a series of numerical values is continued in tabular form, as follows:

1. 100 PER CENT POWER-FACTOR.

$E_0 = 80$ kv. at step-down end of line.

$A = (0.919 + 0.036j)E_0 = 55.1 + 2.2j$ kv.

$D = (0.0144 - 1.168j)E_0 10^{-3} = 0.9 - 70.1j$ amp.

I_0 amp.	B kv.	$E_1 = e_1 + je_2$ $= A + B$	$e_1^2 + e_2^2 = e^2$	e	$\frac{e_2}{e_1} = \tan \epsilon$	$\sum \epsilon$
0	0	55.1 + 2.2j	3036 + 5 = 3041	55.1	+0.040	+2.3
20	1.1 + 2.6j	56.2 + 4.8j	3158 + 23 = 3181	56.4	+0.085	+4.9
40	2.3 + 5.3j	57.4 + 9.5j	3295 + 56 = 3351	57.9	+0.131	+7.5
60	3.4 + 7.9j	58.5 + 10.1j	3422 + 102 = 3524	59.4	+0.173	+9.9
80	4.5 + 10.5j	59.6 + 12.7j	3552 + 161 = 3713	60.9	+0.213	+12.0
100	5.7 + 13.2j	60.8 + 15.4j	3697 + 237 = 3934	62.7	+0.253	+14.2
120	6.8 + 15.8j	61.9 + 18.0j	3832 + 324 = 4156	64.5	+0.291	+16.3

I_0 amp.	C amp.	$I_1 = i_1 + ji_2$ $= C - D$	$i_1^2 + i_2^2 = i^2$	i	$\frac{i_2}{i_1} = \tan \phi$	$\sum \phi$	$\frac{\sum \epsilon}{\sum \phi} = \frac{\sum \epsilon}{\sum \phi}$	Power factor
0	0	-0.7 - 90.1j	4914 + 1 = 4915	70.1	-78	-89.1 + 88.6	0.024	
20	18.4 + 0.7j	17.5 + 70.8j	5013 + 308 = 5319	72.9	+4.04	+78.3 + 71.4	0.332	
40	36.8 + 1.4j	35.9 + 71.5j	5112 + 1289 = 8401	80.0	+1.99	+63.4 + 55.9	0.558	
60	55.1 + 2.2j	54.2 + 72.3j	5227 + 2938 = 8165	90.4	+1.33	+53.1 + 43.2	0.728	
80	73.5 + 2.9j	72.6 + 73.0j	5329 + 5271 = 10600	103.0	+1.055	+45.2 + 33.2	0.837	
100	91.9 + 3.6j	91.0 + 73.9j	8281 + 5432 = 13713	117.1	+0.811	+39.1 + 24.9	0.907	
120	110.3 + 4.3j	109.4 + 74.4j	11969 + 5535 = 17504	132.3	+0.680	34.1 + 17.8	0.952	lead

$e_1 = 60$ kv. at step-up end of line.					
I_0 amp.	Red. Factor, $\frac{e}{60}$	i_0 amp.	e_0 kv.	i_1 amp.	Power-Factor.
0	0.918	0	65.5	76.4	0.024
20	0.940	21.3	63.8	77.5	0.332
40	0.965	41.4	62.1	82.9	0.558
60	0.990	60.6	60.6	91.4	0.728
80	1.015	78.8	59.1	101.5	0.837
100	1.045	95.7	57.5	112.3	0.907
120	1.075	111.7	55.8	122.8	0.952
					lead
Curves of i_0 , e_0 , i_1 , $\cos \theta$, plotted in Fig. 86.					

2. 90 PER CENT POWER-FACTOR, LAG.

$$\cos \theta = 0.9; \quad \sin \theta = \sqrt{1 - 0.9^2} = 0.436;$$

$$I_0 = i_0(\cos \theta - j \sin \theta) = i_0(0.9 - 0.436j);$$

$$\begin{aligned} E_1 &= (0.919 + 0.036j)e_0 + (56.8 + 131.8j)(0.9 - 0.436j)i_0 \\ &= (0.919 + 0.036j)e_0 + (108.5 + 93.8j)i_0 = A + B'; \end{aligned}$$

$$\begin{aligned} I_1 &= (0.919 + 0.036j)(0.9 - 0.436j)i_0 - (0.0144 - 1.168j)e_0 10^{-3} \\ &= (0.843 - 0.366j)i_0 - (0.0144 - 1.168j)e_0 10^{-3} = C' - D, \end{aligned}$$

and now the table is calculated in the same manner as under 1.

Then corresponding tables are calculated, in the same manner, for power-factor, =0.8 and =0.7, respectively, lag, and for power-factor =0.9, 0.8, 0.7, lead; that is, for

$$\begin{aligned} \cos \theta + j \sin \theta &= 0.8 - 0.6j; \\ &0.7 - 0.714j; \\ &0.9 + 0.436j; \\ &0.8 + 0.6j; \\ &0.7 + 0.714j. \end{aligned}$$

Then curves are plotted for all seven values of power-factor, from 0.7 lag to 0.7 lead.

From these curves, for a number of values of i_0 , for instance, $i_0 = 20, 40, 60, 80, 100$, numerical values of i_1 , e_0 , $\cos \theta$, are

taken, and plotted as curves, which, for the same voltage $e_1=60$ at the step-up end, give i_1 , e_0 , and $\cos \theta$, for the same value i_0 , that is, give the regulation of the line at constant current output for varying power-factor.

b. Accuracy of Calculation.

166. Not all engineering calculations require the same degree of accuracy. When calculating the efficiency of a large alternator it may be of importance to determine whether it is 97.7 or 97.8 per cent, that is, an accuracy within one-tenth per cent may be required; in other cases, as for instance, when estimating the voltage which may be produced in an electric circuit by a line disturbance, it may be sufficient to

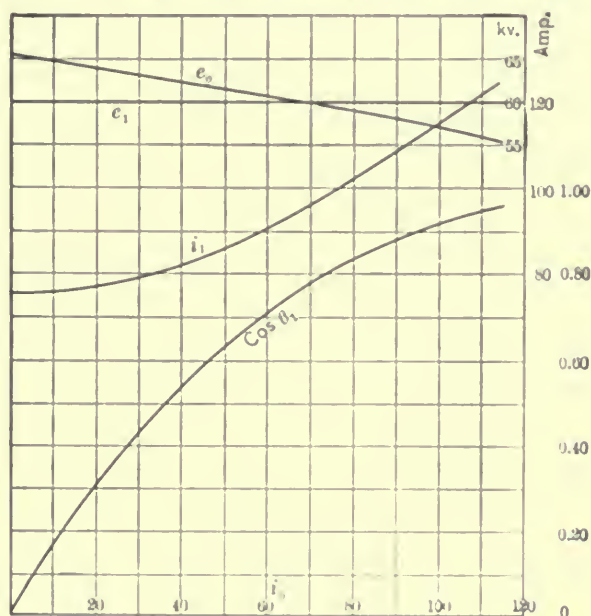


FIG. 102. Transmission Line Characteristics.

determine whether this voltage would be limited to double the normal circuit voltage, or whether it might be 5 or 10 times the normal voltage.

In general, according to the degree of accuracy, engineering calculations may be roughly divided into three classes:

(a) Estimation of the magnitude of an effect; that is, determining approximate numerical values within 25, 50, or 100 per cent. Very frequently such very rough approximation is sufficient, and is all that can be expected or calculated. For instance, when investigating the short-circuit current of an electric generating system, it is of importance to know whether this current is 3 or 4 times normal current, or whether it is 40 to 50 times normal current, but it is immaterial whether it is 45 to 46 or 50 times normal. In studying lightning phenomena, and, in general, abnormal voltages in electric systems, calculating the discharge capacity of lightning arresters, etc., the magnitude of the quantity is often sufficient. In calculating the critical speed of turbine alternators, or the natural period of oscillation of synchronous machines, the same applies, since it is of importance only to see that these speeds are sufficiently remote from the normal operating speed to give no trouble in operation.

(b) Approximate calculation, requiring an accuracy of one or a few per cent only; a large part of engineering calculations fall in this class, especially calculations in the realm of design. Although, frequently, a higher accuracy could be reached in the calculation proper, it would be of no value, since the data on which the calculations are based are susceptible to variations beyond control, due to variation in the material, in the mechanical dimensions, etc.

Thus, for instance, the exciting current of induction motors may vary by several per cent, due to variations of the length of air gap, so small as to be beyond the limits of constructive accuracy, and a calculation exact to a fraction of one per cent, while theoretically possible, thus would be practically useless. The calculation of the ampere-turns required for the shunt field excitation, or for the series field of a direct-current generator needs only moderate exactness, as variations in the magnetic material, in the speed regulation of the driving power, etc., produce differences amounting to several per cent.

(c) Exact engineering calculations, as, for instance, the calculations of the efficiency of apparatus, the regulation of transformers, the characteristic curves of induction motors, etc. These are determined with an accuracy frequently amounting to one-tenth of one per cent and even greater.

Even for most exact engineering calculations, the accuracy of the slide rule is usually sufficient, if intelligently used, that is, used so as to get the greatest accuracy. For accurate calculations, preferably the glass slide should not be used, but the result interpolated by the eye.

Thereby an accuracy within $\frac{1}{4}$ per cent can easily be maintained.

For most engineering calculations, logarithmic tables are sufficient for three decimals, if intelligently used, and as such tables can be contained on a single page, their use makes the calculation very much more expeditious than tables of more decimals. The same applies to trigonometric tables: tables of the trigonometric functions (not their logarithms) of three decimals I find most convenient for most cases, given from degree to degree, and using decimal fractions of the degrees (not minutes and seconds).*

Expedition in engineering calculations thus requires the use of tools of no higher accuracy than required in the result, and such are the slide rules, and the three decimal logarithmic and trigonometric tables. The use of these, however, make it necessary to guard in the calculation against a *loss of accuracy*.

Such loss of accuracy occurs in subtracting or dividing two terms which are nearly equal, in some logarithmic operations, solution of equation, etc., and in such cases either a higher accuracy of calculation must be employed—seven decimal logarithmic tables, etc.—or the operation, which lowers the accuracy, avoided. The latter can usually be done. For instance, in dividing $297 \div 283$ by the slide rule, the proper way is to divide $297 - 283 = 14$ by 283, and add the result to 1.

It is in the methods of calculation that experience and judgment and skill in efficiency of arrangement of numerical calculations is most marked.

167. While the calculations are unsatisfactory, if not carried out with the degree of exactness which is feasible and desirable, it is equally wrong to give numerical values with a number of

* This obviously does not apply to some classes of engineering work, in which a much higher accuracy of trigonometric functions is required, as trigonometric surveying, etc.

ciphers greater than the method or the purpose of the calculation warrants. For instance, if in the design of a direct-current generator, the calculated field ampere-turns are given as 9738, such a numerical value destroys the confidence in the work of the calculator or designer, as it implies an accuracy greater than possible, and thereby shows a lack of judgment.

The number of ciphers in which the result of calculation is given should signify the exactness. In this respect two systems are in use:

(a) Numerical values are given with one more decimal than warranted by the probable error of the result; that is, the decimal before the last is correct, but the last decimal may be wrong by several units. This method is usually employed in astronomy, physics, etc.

(b) Numerical values are given with as many decimals as the accuracy of the calculation warrants; that is, the last decimal is probably correct within half a unit. For instance, an efficiency of 86 per cent means an efficiency between 85.5 and 86.5 per cent; an efficiency of 97.3 per cent means an efficiency between 97.25 and 97.35 per cent, etc. This system is generally used in engineering calculations. To get accuracy of the last decimal of the result, the calculations then must be carried out for one more decimal than given in the result. For instance, when calculating the efficiency by adding the various percentages of losses, data like the following may be given:

Core loss	2.73	per cent
i^2r	1.06	"
Friction	0.93	"
Total	4.72	"
Efficiency	$100 - 4.72 = 95.38$	
Approximately	95.4	"

It is obvious that throughout the same calculation the same degree of accuracy must be observed.

It follows herefrom that the values:

$$2\frac{1}{2}; \quad 2.5; \quad 2.50; \quad 2.500,$$

while mathematically equal, are not equal in their meaning as an engineering result:

- 2.5 means between 2.45 and 2.55;
- 2.50 means between 2.495 and 2.505;
- 2.500 means between 2.4995 and 2.5005;

while $2\frac{1}{2}$ gives no clue to the accuracy of the value.

Thus it is not permissible to add zeros, or drop zeros at the end of numerical values, nor is it permissible, for instance, to replace fractions as $1/16$ by 0.0625, without changing the meaning of the numerical value, as regards its accuracy. This is not always realized, and especially in the reduction of common fractions to decimals an unjustified laxness exists which impairs the reliability of the results. For instance, if in an arc lamp the arc length, for which the mechanism is adjusted, is stated to be 0.8125 inch, such a statement is ridiculous, as no arc lamp mechanism can control for one-tenth as great an accuracy as implied in this numerical value: the value is an unjustified translation from $13/16$ inch.

The principle thus should be adhered to, that all calculations are carried out for one decimal more than the exactness required or feasible, and in the result the last decimal dropped; that is, the result given so that the last decimal is probably correct within half a unit.

c. Intelligibility of Engineering Data and Engineering Reports.

168. In engineering calculations the value of the results mainly depends on the information derived from them, that is, on their intelligibility. To make the numerical results and their meaning as intelligible as possible, it thus is desirable, whenever a series of values are calculated, to carefully arrange them in tables and plot them in a curve or in curves. The latter is necessary, since for most engineers the plotted curve gives a much better conception of the shape and the variation of a quantity than numerical tables.

Even where only a single point is required, as the core loss at full load, or the excitation of an electric generator at rated voltage, it is generally preferable to calculate a few

points near the desired value, so as to get at least a short piece of curve including the desired point.

The main advantage, and foremost purpose of curve plotting thus is to show the shape of the function, and thereby give a clearer conception of it; but for recording numerical values, and deriving numerical values from it, the plotted curve is inferior to the table, due to the limited accuracy possible in a plotted curve, and the further inaccuracy resulting when drawing a curve through the plotted calculated points. To some extent, the numerical values as taken from a plotted curve, depend on the particular kind of curve rule used in plotting the curve.

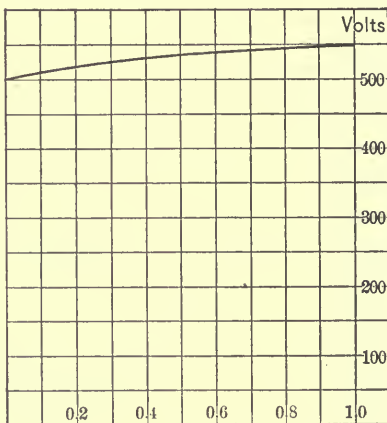


FIG. 103. Compounding Curve.

In general, curves are used for two different purposes, and on the purpose for which the curve is plotted, should depend the method of plotting, as the scale, the zero values, etc.

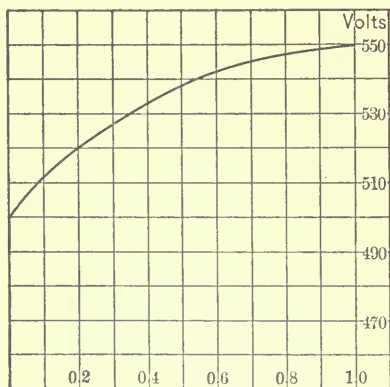


FIG. 104. Compounding Curve.

When curves are used to illustrate the shape of the function, so as to show how much and in what manner a quantity varies as function of another, large divisions of inconspicuous cross-sectioning are desirable, but it is essential that the cross-sectioning should extend to the zero values of the function, even if the numerical values do not extend so far, since otherwise a wrong impression would be con-

ferred. As illustrations are plotted in Figs. 103 and 104, the compounding curve of a direct-current generator. The arrange-

ment in Fig. 103 is correct; it shows the relative variation of voltage as function of the load. Fig. 104, in which the cross-sectioning does not begin at the scale zero, confers the

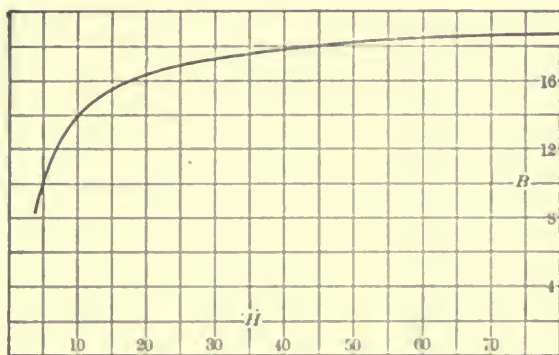


FIG. 105. Curve Plotted to show Characteristic Shape.

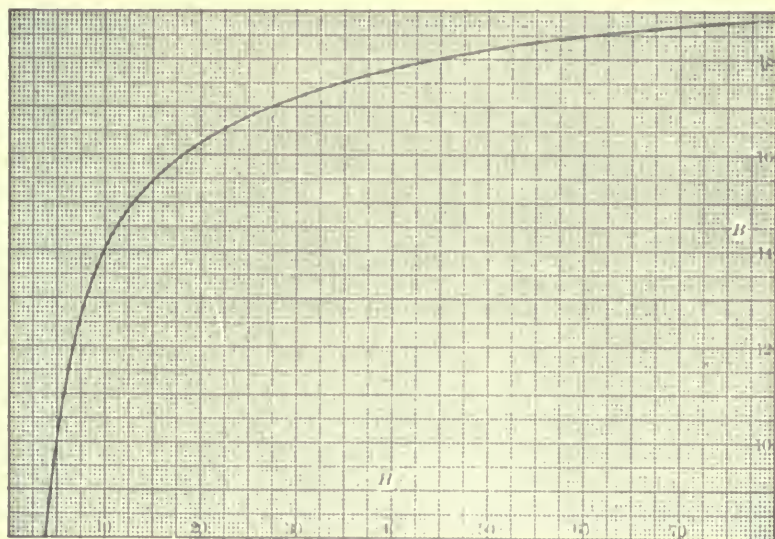


FIG. 106. Curve Plotted for Use as Design Data.

wrong impression that the variation of voltage is far greater than it really is.

When curves are used to record numerical values and derive them from the curve, as, for instance, is commonly the

case with magnetization curves, it is unnecessary to have the zero of the function coincide with the zero of the cross-sectioning, but rather preferable not to have it so, if thereby a better scale of the curve can be secured. It is desirable, however, to use sufficiently small cross-sectioning to make it possible to take numerical values from the curve with good accuracy. This is illustrated by Figs. 105 and 106. Both show the magnetic characteristic of soft steel, for the range above $B = 8000$, in which it is usually employed. Fig. 105 shows the proper way of plotting for showing the shape of the function, Fig. 106 the proper way of plotting for use of the curve to derive numerical values therefrom.

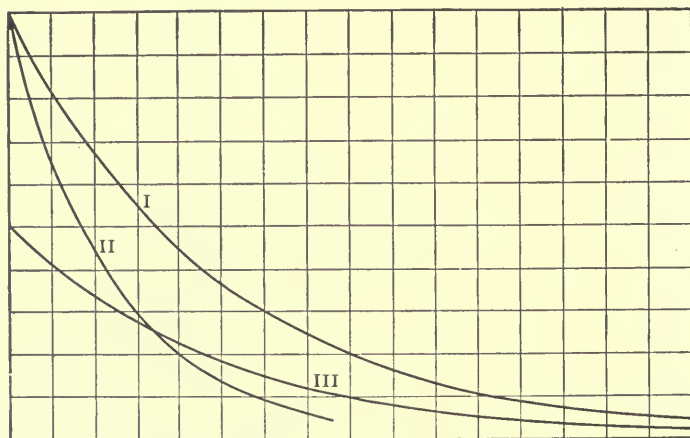


FIG. 107. Same Function Plotted to Different Scales; I is correct.

169. Curves should be plotted in such a manner as to show the quantity which they represent, and its variation, as well as possible. Two features are desirable herefor:

1. To use such a scale that the average slope of the curve, or at least of the more important part of it, does not differ much from 45 deg. Hereby variations of curvature are best shown. To illustrate this, the exponential function $y = e^{-x}$ is plotted in three different scales, as curves I, II, III, in Fig. 107. Curve I has the proper scale.

2. To use such a scale, that the total range of ordinates is not much different from the total range of abscissas. Thus when plotting the power-factor of an induction motor, in Fig. 108, curve I is preferable to curves II or III.

These two requirements frequently are at variance with each other, and then a compromise has to be made between them, that is, such a scale chosen that the total ranges of the two coordinates do not differ much, and at the same time the average slope of the curve is not far from 45 deg. This usually leads to a somewhat rectangular area covered by the curve, as shown, for instance, by curve I, in Fig. 107.

It is obvious that, where the inherent nature of the curve is incompatible with 45 degree slope, this rule does not apply. Such for instance is the case with instrument calibration curves, which inherently are essentially horizontal lines, with curves like the slip of induction motors, etc.

As regards to the magnitude of the scale of plotting, the larger the scale, the plainer obviously is the curve. It must be kept in mind, however, that it would be wrong to use a scale which is materially larger than the accuracy of the values plotted.

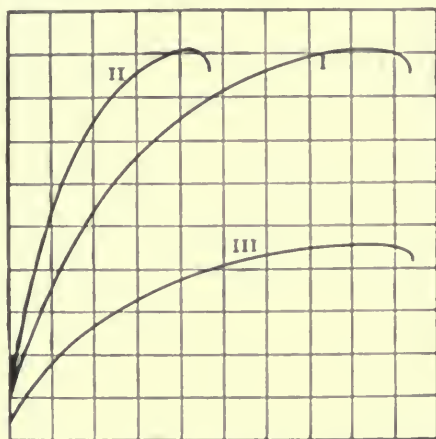


FIG. 108. Same Function Plotted to Different Scales; I is Correct.

Thus for instance, in plotting the calibration curve of an instrument, if the accuracy of the calibration is not greater than .05 per cent, it would be wrong to use .01 per cent as the unit of ordinate scale.

In curve plotting, a scale should be used which is easily read. Hence, only full scale, double scale, and half scale should be used. Triple scale and one-third scale are practically

unreadable, and should therefore never be used. Quadruple scale and quarter scale are difficult to read and therefore undesirable, and are generally unnecessary, since quadruple scale is not much different from half scale with a ten times smaller unit, and quarter scale not much different from double scale of a ten times larger unit.

170. In plotting a curve to show a relation $y=f(x)$, in general x and y should be plotted directly, on ordinary coordinate

paper, but not $\log x$, or y^2 , or logarithmic paper used, etc., as this would not show the shape of the relation $y=f(x)$. Using for instance semi-logarithmic paper, that is, with logarithmic abscissæ and ordinary ordinates, the plotted curve would show the shape of the relation $y=f(\log x)$, etc. The use of logarithmic paper, or the use of y^2 , or $\sqrt{x}$ as coordinate, etc., is justified only where the purpose is to show the relation between y and $\log x$, or between y^2 and x , or between y and $\sqrt{x}$, etc., as is the case when investigating the equation of an empirical curve, or when intending to show some particular feature of the relation $y=f(x)$. Thus for instance when plotting the power p consumed by corona in a high potential transmission line, as function of the line voltage e , by using $\sqrt{p}$ as ordinate, a straight line results. Also where some particular function of one of the coordinates, as $\log x$, gives a more rational relation, it may be used instead of x . Thus for instance in radiation curves, or when expressing velocity as function of wave length or frequency, or expressing attenuation of a wireless wave, etc., the log of wave length or frequency, that is, the geometric scale (as used in the theory of sound, with the octave as unit) is more rational and therefore preferable.

Sometimes the values of a relationship extend over such a wide range as to make it impossible to represent all of them in one curve, and then a number of curves may have to be used, with different scales. In such cases, the logarithmic scale often brings all values within one curve without improperly crowding, and especially where the purpose of curve plotting is not so much to show the shape of the relation, as to record for the purpose of taking numerical values from the curve, the latter arrangement, that is, the use of logarithmic or semi-logarithmic paper may be desirable. Thus the magnetic characteristic of iron is used over a range of field intensities from very few ampere turns per cm. in transformers, to thousands of ampere turns, in tooth densities of railway motors, and the magnetic characteristic thus is either represented by three curves with different scales, of ratios $1 \div 10 \div 100$, as shown in Fig. 109, or the log of field intensity used as abscissæ, that is, semi-logarithmic paper, with logarithmic scale as abscissæ, and regular scale as ordinates, as shown in Fig. 110.

It must be realized that the logarithmic or geometrical scale

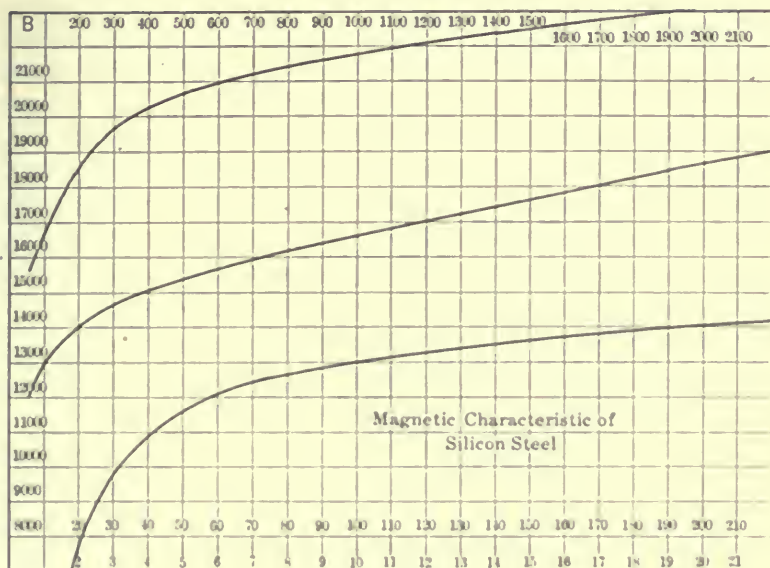


FIG. 109.

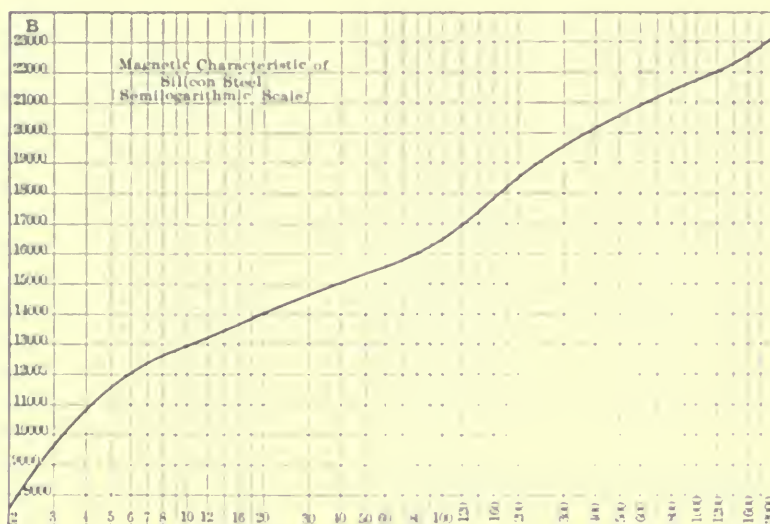


FIG. 110.

—in which equal divisions represent *not* equal values of the quantity, but equal fractions of the quantity—is somewhat less easy to read than common scale. However, as it is the same scale as the slide rule, this is not a serious objection.

A disadvantage of the logarithmic scale is that it cannot extend down to zero, and relations in which the entire range down to zero requires consideration, thus are not well suited for the use of logarithmic scale.

171. Any engineering calculation on which it is worth while to devote any time, is worth being recorded with sufficient completeness to be generally intelligible. Very often in making calculations the data on which the calculation is based, the subject and the purpose of the calculation are given incompletely or not at all, since they are familiar to the calculator at the time of calculation. The calculation thus would be unintelligible to any other engineer, and usually becomes unintelligible even to the calculator in a few weeks.

In addition to the name and the date, all calculations should be accompanied by a complete record of the object and purpose of the calculation, the apparatus, the assumptions made, the data used, reference to other calculations or data employed, etc., in short, they should include all the information required to make the calculation intelligible to another engineer without further information besides that contained in the calculations, or in the references given therein. The small amount of time and work required to do this is negligible compared with the increased utility of the calculation.

Tables and curves belonging to the calculation should in the same way be completely identified with it and contain sufficient data to be intelligible.

171A. Engineering investigations evidently are of no value, unless they can be communicated to those to whom they are of interest. Thus the engineering report is an essential and important part of the work. If therefore occasionally an engineer or scientist is met, who is so much interested in the investigating work, that he hates to “waste” the time of making proper and complete reports, this is a very foolish attitude, since in general it destroys the value of the work.

As practically every engineering investigation is of interest and importance to different classes of people, as a rule not one,

but several reports must be written to make the most use of the work; the scientific record of the research would be of no more value to the financial interests considering the industrial development of the work than the report to the financial or administrative body would be of value to the scientist, who considers repeating and continuing the investigation.

In general thus three classes of engineering reports can be distinguished, and all three reports should be made with every engineering investigation, to get best use of it.

(a) The scientific record of the investigation. This must be so complete as to enable another investigator to completely check up, repeat and further extend the investigation. It thus must contain the original observations, the method of work, apparatus and facilities, calibrations, information on the limits of accuracy and reliability, sources of error, methods of calculation, etc., etc.

It thus is a lengthy report, and as such will be read by very few, if any, except other competent investigators, but is necessary as the record of the work, since without such report, the work would be lost, as the conclusions and results could not be checked up if required.

This report appeals only to men of the same character as the one who made the investigation, and is essentially for record and file.

(b) The general engineering report. It should be very much shorter than the scientific report, should be essentially of the nature of a syllabus thereof, avoid as much as possible complex mathematical and theoretical considerations, but give all the engineering results of the investigation, in as plain language as possible. It would be addressed to administrative engineers, that is, men who as engineers are capable of understanding the engineering results and discussion, but have neither time nor familiarity to follow in detail through the investigation, and are not interested in such things as the original readings, the discussion of methods, accuracy, etc., but are interested only in the results.

This is the report which would be read by most of the men interested in the matter. It would in general be the form in which the investigation is communicated to engineering societies as paper, with the scientific report relegated into an appendix of the paper.

(c) The general report. This should give the results, that is, explain what the matter is about, in plain and practically non-technical language, addressed to laymen, that is, non-engineers. In other words, it should be understood by any intelligent non-technical man.

Such general report would be materially shorter than the general engineering report, as it would omit all details, and merely deal with the general problem, purpose and solution.

In general, it is advisable to combine all three reports, by having the scientific record preceded by the general engineering report, and the latter preceded by the general report. Roughly, the general report would usually have a length of 20 to 40 per cent of the general engineering report, the latter a length of 10 to 25 per cent of the complete scientific record.

The bearing of the three classes of reports may be understood by illustration on an investigation which appears of commercial utility, and therefore is submitted for industrial development to a manufacturing corporation; the financial and general administrative powers of the corporations, to whom the investigation is submitted, would read the general report and if the matter appears to them of sufficient interest for further consideration, refer it to the engineering department. The general report thus must be written for, and intelligible to the non-engineering administrative heads of the organization. The administrative engineers of the engineering department then peruse the general engineering report, and this report must be an engineering report, but general and not require the knowledge of the specialist in the particular field. If then the conclusion derived by the administrative engineers from the reading of the general engineering report is to the effect that the matter is worth further consideration, then they refer it to the specialists in the field covered by the investigation, and to the latter finally the scientific record of the investigation appeals and is studied in making final report on the work.

Inversely, where nothing but a lengthy scientific report is submitted, as a rule it will be referred to the engineering department, and the general engineer, even if he could wade through the lengthy report, rarely has immediately time to do so, thus lays it aside to study sometime at his leisure—and very often this time never comes, and the entire matter drops, for lack of proper representation.

Thus it is of the utmost importance for the engineer and the scientist, to be able to present the results of his work not only by elaborate and lengthy scientific report, but also by report of moderate length, intelligible without difficulty to the general engineer, and by short statement intelligible to the non-engineer.

d. Reliability of Numerical Calculations.

172. The most important and essential requirement of numerical engineering calculations is their absolute reliability. When making a calculation, the most brilliant ability, theoretical knowledge and practical experience of an engineer are made useless, and even worse than useless, by a single error in an important calculation.

Reliability of the numerical calculation is of vastly greater importance in engineering than in any other field. In pure mathematics an error in the numerical calculation of an example which illustrates a general proposition, does not detract from the interest and value of the latter, which is the main purpose; in physics, the general law which is the subject of the investigation remains true, and the investigation of interest and use, even if in the numerical illustration of the law an error is made. With the most brilliant engineering design, however, if in the numerical calculation of a single structural member an error has been made, and its strength thereby calculated wrong, the rotor of the machine flies to pieces by centrifugal forces, or the bridge collapses, and with it the reputation of the engineer. The essential difference between engineering and purely scientific calculations is the rapid check on the correctness of the calculation, which is usually afforded by the performance of the calculated structure—but too late to correct errors.

Thus rapidity of calculation, while by itself useful, is of no value whatever compared with reliability—that is, correctness.

One of the first and most important requirements to secure reliability is neatness and care in the execution of the calculation. If the calculation is made on any kind of a sheet of paper, with lead pencil, with frequent striking out and correcting of figures, etc., it is practically hopeless to expect correct results from any more extensive calculations. Thus the work

should be done with pen and ink, on white ruled paper; if changes have to be made, they should preferably be made by erasing, and not by striking out. In general, the appearance of the work is one of the best indications of its reliability. The arrangement in tabular form, where a series of values are calculated, offers considerable assistance in improving the reliability.

173. Essential in all extensive calculations is a complete system of checking the results, to insure correctness.

One way is to have the same calculation made independently by two different calculators, and then compare the results. Another way is to have a few points of the calculation checked by somebody else. Neither way is satisfactory, as it is not always possible for an engineer to have the assistance of another engineer to check his work, and besides this, an engineer should and must be able to make numerical calculations so that he can absolutely rely on their correctness without somebody else assisting him.

In any more important calculations every operation thus should be performed twice, preferably in a different manner. Thus, when multiplying or dividing by the slide rule, the multiplication or division should be repeated mentally, approximately, as check; when adding a column of figures, it should be added first downward, then as check upward, etc.

Where an exact calculation is required, first the magnitude of the quantity should be estimated, if not already known, then an approximate calculation made, which can frequently be done mentally, and then the exact calculation; or, inversely, after the exact calculation, the result may be checked by an approximate mental calculation.

Where a series of values is to be calculated, it is advisable first to calculate a few individual points, and then, entirely independently, calculate in tabular form the series of values, and then use the previously calculated values as check. Or, inversely, after calculating the series of values a few points should independently be calculated as check.

When a series of values is calculated, it is usually easier to secure reliability than when calculating a single value, since in the former case the different values check each other. Therefore it is always advisable to calculate a number of values, that is, a short curve branch, even if only a single point is

required. After calculating a series of values, they are plotted as a curve to see whether they give a smooth curve. If the entire curve is irregular, the calculation should be thrown away, and the entire work done anew, and if this happens repeatedly with the same calculator, the calculator is advised to find another position more in agreement with his mental capacity. If a single point of the curve appears irregular, this points to an error in its calculation, and the calculation of the point is checked; if the error is not found, this point is calculated entirely separately, since it is much more difficult to find an error which has been made than it is to avoid making an error.

174. Some of the most frequent numerical errors are:

1. The decimal error, that is, a misplaced decimal point. This should not be possible in the final result, since the magnitude of the latter should by judgment or approximate calculation be known sufficiently to exclude a mistake by a factor 10. However, under a square root or higher root, in the exponent of a decreasing exponential function, etc., a decimal error may occur without affecting the result so much as to be immediately noticed. The same is the case if the decimal error occurs in a term which is relatively small compared with the other terms, and thereby does not affect the result very much. For instance, in the calculation of the induction motor characteristics, the quantity $r_1^2 + s^2 x_1^2$ appears, and for small values of the slip s , the second term $s^2 x_1^2$ is small compared with r_1^2 , so that a decimal error in it would affect the total value sufficiently to make it seriously wrong, but not sufficiently to be obvious.

2. Omission of the factor or divisor 2.

3. Error in the sign, that is, using the plus sign instead of the minus sign, and inversely. Here again, the danger is especially great, if the quantity on which the wrong sign is used combines with a larger quantity, and so does not affect the result sufficiently to become obvious.

4. Omitting entire terms of smaller magnitude, etc.

APPENDIX A.

NOTES ON THE THEORY OF FUNCTIONS.

A. General Functions.

175. The most general algebraic expression of powers of x and y ,

$$\begin{aligned} F(x,y) = & (a_{00} + a_{01}x + a_{02}x^2 + \dots) + (a_{10} + a_{11}x + a_{12}x^2 + \dots)y \\ & + (a_{20} + a_{21}x + a_{22}x^2 + \dots)y^2 + \dots \\ & + (a_{n0} + a_{n1}x + a_{n2}x^2 + \dots)y^n = 0, \quad . \quad . \quad . \quad . \quad (1) \end{aligned}$$

is the *implicit analytic function*. It relates y and x so that to every value of x there correspond n values of y , and to every value of y there correspond m values of x , if m is the exponent of the highest power of x in (1).

Assuming expression (1) solved for y (which usually cannot be carried out in final form, as it requires the solution of an equation of the n th order in y , with coefficients which are expressions of x), the *explicit analytic function*,

$$y = f(x), \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (2)$$

is obtained. Inversely, solving the implicit function (1) for x , that is, from the explicit function (2), expressing x as function of y , gives the *reverse function* of (2); that is

$$x = f_1(y). \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (3)$$

In the general algebraic function, in its implicit form (1), or the explicit form (2), or the reverse function (3), x and y are assumed as general numbers; that is, as complex quantities; thus,

$$\left. \begin{aligned} x &= x_1 + jx_2; \\ y &= y_1 + jy_2; \end{aligned} \right\} . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (4)$$

and likewise are the coefficients $a_{00}, a_{01} \dots a_{nm}$.

If all the coefficients a are real, and x is real, the corresponding n values of y are either real, or pairs of conjugate complex imaginary quantities: $y_1 + jy_2$ and $y_1 - jy_2$.

176. For $n=1$, the implicit function (1), solved for y , gives the *rational function*,

$$y = \frac{a_{00} + a_{01}x + a_{02}x^2 + \dots}{a_{10} + a_{11}x + a_{12}x^2 + \dots}, \quad \dots \quad (5)$$

and if in this function (5) the denominator contains no x , the *integer function*,

$$y = a_0 + a_1x + a_2x^2 + \dots + a_mx^m, \quad \dots \quad (6)$$

is obtained.

For $n=2$, the implicit function (1) can be solved for y as a quadratic equation, and thereby gives

$$y = \frac{-(a_{10} + a_{11}x + a_{12}x^2 + \dots) \pm \sqrt{(a_{10} + a_{11}x + a_{12}x^2 + \dots)^2 - 4(a_{00} + a_{01}x + a_{02}x^2 + \dots)(a_{20} + a_{21}x + a_{22}x^2 + \dots)}}{2(a_{20} + a_{21}x + a_{22}x^2 + \dots)}; \quad (7)$$

that is, the explicit form (2) of equation (1) contains in this case a square root.

For $n > 2$, the explicit form $y = f(x)$ either becomes very complicated, for $n=3$ and $n=4$, or cannot be produced in finite form, as it requires the solution of an equation of more than the fourth order. Nevertheless, y is still a function of x , and can as such be calculated by approximation, etc.

To find the value y_1 , which by function (1) corresponds to $x = x_1$, Taylor's theorem offers a rapid approximation. Substituting x_1 in function (1) gives an expression which is of the n th order in y , thus: $F(x_1y)$, and the problem now is to find a value y_1 , which makes $F(x_1, y_1) = 0$.

However,

$$F(x_1, y_1) = F(x_1, y) + h \frac{dF(x_1, y)}{dy} + \frac{h^2}{2} \frac{d^2F(x_1, y)}{d^2y^2} + \dots \quad (8)$$

where $h = y_1 - y$ is the difference between the correct value y_1 and any chosen value y .

Neglecting the higher orders of the small quantity h , in (8), and considering that $F(x_1, y_1) = 0$, gives

$$h = -\frac{F(x_1, y)}{\frac{dF(x_1, y)}{dy}}, \quad \dots \dots \dots (9)$$

and herefrom is obtained $y_1 = y + h$, as first approximation. Using this value of y_1 as y in (9) gives a second approximation, which usually is sufficiently close.

177. New functions are defined by the integrals of the analytic functions (1) or (2), and by their reverse functions. They are called *Abelian integrals* and *Abelian functions*.

Thus in the most general case (1), the explicit function corresponding to (1) being

$$y = f(x), \quad \dots \dots \dots (2)$$

the integral,

$$z = \int f(x) dx,$$

then is the general Abelian integral, and its reverse function,

$$x = \phi(z),$$

the general Abelian function.

(a) In the case, $n = 1$, function (2) gives the rational function (5), and its special case, the integer function (6).

Function (6) can be integrated by powers of x . (5) can be resolved into partial fractions, and thereby leads to integrals of the following forms:

$$\left. \begin{aligned} (1) \quad & \int x^m dx; \\ (2) \quad & \int \frac{dx}{x-a}; \\ (3) \quad & \int \frac{dx}{(x-a)^m}; \\ (4) \quad & \int \frac{dx}{1+x^2}; \end{aligned} \right\} \dots \dots \dots (10)$$

Integrals (10), (1), and (3) integrated give rational functions, (10), (2) gives the logarithmic function $\log(x-a)$, and (10), (4) the arc function $\arctan x$.

As the arc functions are logarithmic functions with complex imaginary argument, this case of the integral of the rational function thus leads to the *logarithmic function*, or the logarithmic integral, which in its simplest form is

$$z = \int \frac{dx}{x} = \log x, \quad \dots \quad (11)$$

and gives as its reverse function the *exponential function*,

$$x = e^z. \quad \dots \quad (12)$$

It is expressed by the infinite series,

$$e^z = 1 + z + \frac{z^2}{2} + \frac{z^3}{3} + \frac{z^4}{4} + \dots \quad \dots \quad (13)$$

as seen in Chapter II, paragraph 53.

178. b. In the case, $n=2$, function (2) appears as the expression (7), which contains a square root of some power of x . Its first part is a rational function, and as such has already been discussed in *a*. There remains thus the integral function,

$$z = \int \frac{\sqrt{b_0 + b_1x + b_2x^2 + \dots + b_px^p}}{c_0 + c_1x + c_2x^2 + \dots} dx. \quad \dots \quad (14)$$

This expression (14) leads to a series of important functions.

(1) For $p=1$ or 2,

$$z = \int \frac{\sqrt{b_0 + b_1x + b_2x^2}}{c_0 + c_1x + c_2x^2 + \dots} dx. \quad \dots \quad (15)$$

By substitution, resolution into partial fractions, and separation of rational functions, this integral (11) can be reduced to the standard form,

$$z = \int \frac{dx}{\sqrt{1 \mp x^2}}. \quad \dots \quad (16)$$

In the case of the minus sign, this gives

$$z = \int \frac{dx}{\sqrt{1 - x^2}} = \arcsin x, \quad \dots \quad (17)$$

and as reverse functions thereof, there are obtained the *trigonometric functions*.

$$\left. \begin{aligned} x &= \sin z, \\ \sqrt{1-x^2} &= \cos z. \end{aligned} \right\} \cdot \cdot \cdot \cdot \cdot \cdot \quad (18)$$

In the case of the plus sign, integral (16) gives

$$z = \int \frac{dx}{\sqrt{1+x^2}} = -\log\{\sqrt{1+x^2}-x\} = \text{arc sinh } x, \quad (19)$$

and reverse functions thereof are the *hyperbolic functions*,

$$\left. \begin{aligned} x &= \frac{\epsilon^{+z} - \epsilon^{-z}}{2} = \sinh z; \\ \sqrt{1+x^2} &= \frac{\epsilon^{+z} + \epsilon^{-z}}{2} = \cosh z. \end{aligned} \right\} \cdot \cdot \cdot \cdot \cdot \cdot \quad (20)$$

The trigonometric functions are expressed by the series:

$$\left. \begin{aligned} \sin z &= z - \frac{z^3}{3} + \frac{z^5}{5} - \frac{z^7}{7} + \dots; \\ \cos z &= 1 - \frac{z^2}{2} + \frac{z^4}{4} - \frac{z^6}{6} + \dots, \end{aligned} \right\} \cdot \cdot \cdot \cdot \cdot \cdot \quad (21)$$

as seen in Chapter II, paragraph 58.

The hyperbolic functions, by substituting for ϵ^{+z} and ϵ^{-z} the series (13), can be expressed by the series:

$$\left. \begin{aligned} \sinh z &= z + \frac{z^3}{3} + \frac{z^5}{5} + \frac{z^7}{7} + \dots; \\ \cosh z &= 1 + \frac{z^2}{2} + \frac{z^4}{4} + \frac{z^6}{6} + \dots \end{aligned} \right\} \cdot \cdot \cdot \cdot \cdot \cdot \quad (22)$$

179. In the next case, $p=3$ or 4,

$$z = \int \frac{\sqrt{b_0 + b_1x + b_2x^2 + b_3x^3 + b_4x^4}}{c_0 + c_1x + c_2x^2 + \dots} dx, \quad (23)$$

already leads beyond the elementary functions, that is, (23) cannot be integrated by rational, logarithmic or arc functions,

but gives a new class of functions, the *elliptic integrals*, and their reverse functions, the *elliptic functions*, so called, because they bear to the ellipse a relation similar to that, which the trigonometric functions bear to the circle and the hyperbolic functions to the equilateral hyperbola.

The integral (23) can be resolved into elementary functions, and the three classes of elliptic integrals:

$$\left. \begin{aligned} u &= \int \frac{dx}{\sqrt{x(1-x)(1-c^2x)}}; \\ u_1 &= \int \frac{x dx}{\sqrt{x(1-x)(1-c^2x)}}; \\ u_2 &= \int \frac{dx}{(x-b)\sqrt{x(1-x)(1-c^2x)}}. \end{aligned} \right\} \dots \dots (24)$$

(These three classes of integrals may be expressed in several different forms.)

The reverse functions of the elliptic integrals are given by the elliptic functions:

$$\left. \begin{aligned} \sqrt{x} &= \sin am(u, c); \\ \sqrt{1-x} &= \cos am(u, c); \\ \sqrt{1-c^2x} &= \Delta am(u, c); \end{aligned} \right\} \dots \dots (25)$$

known, respectively, as sine-amplitude, cosine-amplitude, delta-amplitude.

Elliptic functions are in some respects similar to trigonometric functions, as is seen, but they are more general, depending, as they do, not only on the variable x , but also on the constant c . They have the interesting property of being *doubly periodic*. The trigonometric functions are periodic, with the periodicity 2π , that is, repeat the same values after every change of the angle by 2π . The elliptic functions have two periods p_1 and p_2 , that is,

$$\sin am(u + np_1 + mp_2, c) = \sin am(u, c), \text{ etc.}; \dots (26)$$

hence, increasing the variable u by any multiple of either period p_1 and p_2 , repeats the same values.

The two periods are given by the equations,

$$\left. \begin{aligned} p_1 &= \int_0^1 \frac{dx}{2\sqrt{x(1-x)(1-c^2x)}}; \\ p_2 &= \int_1^{\frac{1}{c^2}} \frac{dx}{2\sqrt{x(1-x)(1-c^2x)}}. \end{aligned} \right\} \dots \dots \dots (27)$$

180. Elliptic functions can be expressed as ratios of two infinite series, and these series, which form the numerator and the denominator of the elliptic function, are called *theta functions* and expressed by the symbol θ , thus

$$\left. \begin{aligned} \sin am(u, c) &= \frac{1}{\sqrt{c}} \frac{\theta_1\left(\frac{\pi u}{2p_1}\right)}{\theta_0\left(\frac{\pi u}{2p_1}\right)}; \\ \cos am(u, c) &= \sqrt[4]{\frac{c^2}{1-c^2}} \frac{\theta_2\left(\frac{\pi u}{2p_1}\right)}{\theta_0\left(\frac{\pi u}{2p_1}\right)}; \\ \Delta am(u, c) &= \sqrt[4]{1-c^2} \frac{\theta_3\left(\frac{\pi u}{2p_1}\right)}{\theta_0\left(\frac{\pi u}{2p_1}\right)}, \end{aligned} \right\} \dots \dots \dots (28)$$

and the four θ functions may be expressed by the series:

$$\left. \begin{aligned} \theta_0(x) &= 1 - 2q \cos 2x + 2q^4 \cos 4x - 2q^9 \cos 6x + \dots; \\ \theta_1(x) &= 2q^{1/4} \sin x - 2q^{9/4} \sin 3x + 2q^{\frac{25}{4}} \sin 5x - \dots; \\ \theta_2(x) &= 2q^{1/4} \cos x + 2q^{9/4} \cos 3x + 2q^{\frac{25}{4}} \cos 5x + \dots; \\ \theta_3(x) &= 1 + 2q \cos 2x + 2q^4 \cos 4x + 2q^9 \cos 6x + \dots, \end{aligned} \right\} \dots \dots \dots (29)$$

where

$$q = \varepsilon^a \quad \text{and} \quad a = j\pi \frac{p_2}{p_1}. \quad \dots \dots \dots (30)$$

In the case of integral function (14), where $p > 4$, similar integrals and their reverse functions appear, more complex

than the elliptic functions, and of a greater number of periodicities. They are called *hyperelliptic integrals* and *hyperelliptic functions*, and the latter are again expressed by means of auxiliary functions, the *hyperelliptic θ functions*.

181. Many problems of physics and of engineering lead to elliptic functions, and these functions thus are of considerable importance. For instance, the motion of the pendulum is expressed by elliptic functions of time, and its period thereby is a function of the amplitude, increasing with increasing amplitude; that is, in the so-called "second pendulum," the time of one swing is not constant and equal to one second, but only approximately so. This approximation is very close, as long as the amplitude of the swing is very small and constant, but if the amplitude of the swing of the pendulum varies and reaches large values, the time of the swing, or the period of the pendulum, can no longer be assumed as constant and an exact calculation of the motion of the pendulum by elliptic functions becomes necessary.

In electrical engineering, one has frequently to deal with oscillations similar to those of the pendulum, for instance, in the hunting or surging of synchronous machines. In general, the frequency of oscillation is assumed as constant, but where, as in cumulative hunting of synchronous machines, the amplitude of the swing reaches large values, an appreciable change of the period must be expected, and where the hunting is a resonance effect with some other periodic motion, as the engine rotation, the change of frequency with increase of amplitude of the oscillation breaks the complete resonance and thereby tends to limit the amplitude of the swing.

182. As example of the application of elliptic integrals, may be considered the determination of the length of the arc of an ellipse.

Let the ellipse of equation

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1, \quad . \quad . \quad . \quad . \quad . \quad . \quad (31)$$

be represented in Fig. 93, with the circumscribed circle,

$$x^2 + y^2 = a^2. \quad . \quad . \quad . \quad . \quad . \quad . \quad (32)$$

To every point $P=x, y$ of the ellipse then corresponds a point $P_1=x, y_1$ on the circle, which has the same abscissa x , and an angle $\theta=AOP_1$.

The arc of the ellipse, from A to P , then is given by the integral,

$$L=a \int_0^z \frac{(1-c^2z)dz}{2\sqrt{z(1-z)(1-c^2z)}}, \quad (33)$$

where

$$z=\sin^2 \theta = \left(\frac{x}{a}\right)^2 \quad \text{and} \quad c=\frac{\sqrt{a^2-b^2}}{a}, \quad . . . (34)$$

is the eccentricity of the ellipse.

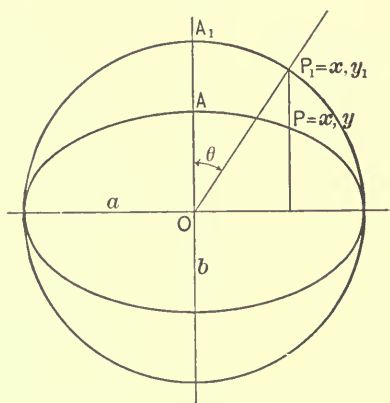


FIG. 93. Rectification of Ellipse.

Thus the problem leads to an elliptic integral of the first and of the second class.

For more complete discussion of the elliptic integrals and the elliptic functions, reference must be made to the text-books of mathematics.

B. Special Functions.

183. Numerous special functions have been derived by the exigencies of mathematical problems, mainly of astronomy, but in the latter decades also of physics and of engineering. Some of them have already been discussed as special cases of the general Abelian integral and its reverse function, as the exponential, trigonometric, hyperbolic, etc., functions.

Functions may be represented by an infinite series of terms; that is, as a sum of an infinite number of terms, which progressively decrease, that is, approach zero. The denotation of the terms is commonly represented by the summation sign Σ .

Thus the exponential functions may be written, when defining,

$$\underline{0}=1; \quad \underline{n}=1 \times 2 \times 3 \times 4 \times \dots \times n,$$

as

$$e^x = 1 + x + \frac{x^2}{\underline{2}} + \frac{x^3}{\underline{3}} + \dots = \sum_0^n \frac{x^n}{\underline{n}}, \quad \dots \quad (35)$$

which means, that terms $\frac{x^n}{\underline{n}}$ are to be added for all values of n from $n=0$ to $n=\infty$.

The trigonometric and hyperbolic functions may be written in the form:

$$\sin x = x - \frac{x^3}{\underline{3}} + \frac{x^5}{\underline{5}} - \frac{x^7}{\underline{7}} + \dots = \sum_0^n (-1)^n \frac{x^{2n+1}}{\underline{2n+1}}; \quad \dots \quad (36)$$

$$\cos x = 1 - \frac{x^2}{\underline{2}} + \frac{x^4}{\underline{4}} - \frac{x^6}{\underline{6}} + \dots = \sum_0^n (-1)^n \frac{x^{2n}}{\underline{2n}}; \quad \dots \quad (37)$$

$$\sinh x = x + \frac{x^3}{\underline{3}} + \frac{x^5}{\underline{5}} + \frac{x^7}{\underline{7}} + \dots = \sum_0^n \frac{x^{2n+1}}{\underline{2n+1}}; \quad \dots \quad (38)$$

$$\cosh x = 1 + \frac{x^2}{\underline{2}} + \frac{x^4}{\underline{4}} + \frac{x^6}{\underline{6}} + \dots = \sum_0^n \frac{x^{2n}}{\underline{2n}}. \quad \dots \quad (39)$$

Functions also may be expressed by a series of factors; that is, as a product of an infinite series of factors, which progressively approach unity. The product series is commonly represented by the symbol $\prod$.

Thus, for instance, the sine function can be expressed in the form,

$$\sin x = x \left(1 - \frac{x^2}{\pi^2}\right) \left(1 - \frac{x^2}{4\pi^2}\right) \left(1 - \frac{x^2}{9\pi^2}\right) \dots = x \prod_1^n \left(1 - \frac{x^2}{n^2\pi^2}\right). \quad (40)$$

184. Integration of known functions frequently leads to new functions. Thus from the general algebraic functions were

derived the Abelian functions. In physics and in engineering, integration of special functions in this manner frequently leads to new special functions.

For instance, in the study of the propagation through space, of the magnetic field of a conductor, in wireless telegraphy, lightning protection, etc., we get new functions. If $i=f(t)$ is the current in the conductor, as function of the time t , at a distance x from the conductor the magnetic field lags by the time $t_1=\frac{x}{S}$, where S is the speed of propagation (velocity of light). Since the field intensity decreases inversely proportional to the distance x , it thus is proportional to

$$y = \frac{f\left(t - \frac{x}{S}\right)}{x}; \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (41)$$

and the total magnetic flux then is

$$\begin{aligned} z &= \int y dx \\ &= \int \frac{f\left(t - \frac{x}{S}\right)}{x} dx. \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (42) \end{aligned}$$

If the current is an alternating current, that is, $f(t)$ a trigonometric function of time, equation (42) leads to the functions,

$$\left. \begin{aligned} u &= \int \frac{\sin x}{x} dx; \\ v &= \int \frac{\cos x}{x} dx. \end{aligned} \right\} . \quad . \quad . \quad . \quad . \quad . \quad . \quad (43)$$

If the current is a direct current, rising as exponential function of the time, equation (42) leads to the function,

$$w = \int \frac{e^x}{x} dx. \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad (44)$$

Substituting in (43) and (44), for $\sin x$, $\cos x$, e^x their infinite series (21) and (13), and then integrating, gives the following:

$$\left. \begin{aligned} \int \frac{\sin x}{x} dx &= x - \frac{x^3}{3 \cdot 3} + \frac{x^5}{5 \cdot 5} - \frac{x^7}{7 \cdot 7} + \dots; \\ \int \frac{\cos x}{x} dx &= \log x - \frac{x^2}{2 \cdot 2} + \frac{x^4}{4 \cdot 4} - \frac{x^6}{6 \cdot 6} + \dots; \\ \int \frac{e^x}{x} dx &= \log x + x + \frac{x^2}{2 \cdot 2} + \frac{x^3}{3 \cdot 3} + \dots \end{aligned} \right\} \quad (45)$$

For further discussion and tables of these functions see "Theory and Calculation of Transient Electric Phenomena and Oscillations," Section III, Chapter VIII, and Appendix.

185. If $y=f(x)$ is a function of x , and $z = \int f(x)dx = \phi(x)$ its integral, the definite integral, $Z = \int_a^b f(x)dx$, is no longer a function of x but a constant,

$$Z = \phi(b) - \phi(a).$$

For instance, if $y=c(x-n)^2$, then

$$z = \int c(x-n)^2 dx = \frac{c(x-n)^3}{3},$$

and the definite integral is

$$Z = \int_a^b c(x-n)^2 dx = \frac{c}{3} \{ (b-n)^3 - (a-n)^3 \}.$$

This definite integral does not contain x , but it contains all the constants of the function $f(x)$, thus is a function of these constants c and n , as it varies with a variation of these constants.

In this manner new functions may be derived by definite integrals.

Thus, if

$$y = f(x, u, v \dots) \quad (46)$$

is a function of x , containing the constants $u, v \dots$

The definite integral,

$$Z = \int_a^b f(x, u, v \dots) dx, \quad . \quad . \quad . \quad . \quad . \quad (47)$$

is not a function of x , but still is a function of $u, v \dots$, and may be a new function.

186. For instance, let

$$y = e^{-x} x^{u-1}; \quad . \quad . \quad . \quad . \quad . \quad (48)$$

then the integral,

$$f(u) = \int_0^\infty e^{-x} x^{u-1} dx, \quad . \quad . \quad . \quad . \quad . \quad (49)$$

is a new function of u , called the *gamma function*.

Some properties of this function may be derived by partial integration, thus:

$$\Gamma(u+1) = u\Gamma(u); \quad . \quad . \quad . \quad . \quad . \quad . \quad (50)$$

if n is an integer number,

$$\Gamma(u) = (u-1)(u-2) \dots (u-n)\Gamma(u-n), \quad . \quad . \quad (51)$$

and since

$$\Gamma(1) = 1, \quad . \quad . \quad . \quad . \quad . \quad . \quad (52)$$

if u is an integer number, then,

$$\Gamma(u) = \underline{u-1}. \quad . \quad . \quad . \quad . \quad . \quad . \quad (53)$$

C. Exponential, Trigonometric and Hyperbolic Functions.

(a) FUNCTIONS OF REAL VARIABLES.

187. The exponential, trigonometric, and hyperbolic functions are defined as the reverse functions of the integrals,

$$a. \quad u = \int \frac{dx}{x} = \log x, \quad . \quad . \quad . \quad . \quad . \quad . \quad (54)$$

$$\text{and:} \quad x = e^u \quad . \quad . \quad . \quad . \quad . \quad . \quad (55)$$

$$b. \quad u = \int \frac{dx}{\sqrt{1-x^2}} = \arcsin x; \quad . \quad . \quad . \quad . \quad . \quad . \quad (56)$$

and: $x = \sin u,$ (57)

$$\sqrt{1-x^2} = \cos u, \quad (58)$$

c. $u = \int \frac{dx}{\sqrt{1+x^2}} = -\log\{\sqrt{1+x^2}-x\}; \quad (59)$

and $x = \frac{e^u - e^{-u}}{2} = \sinh u; \quad (60)$

$$\sqrt{1+x^2} = \frac{e^u + e^{-u}}{2} = \cosh u. \quad (61)$$

From (57) and (58) it follows that

$$\sin^2 u + \cos^2 u = 1. \quad (62)$$

From (60) and (61) it follows that

$$\cos^2 hu - \sin^2 hu = 1. \quad (63)$$

Substituting $(-x)$ for x in (56), gives $(-u)$ instead of u , and therefrom,

$$\sin(-u) = -\sin u \quad (64)$$

Substituting $(-u)$ for u in (60), reverses the sign of x , that is,

$$\sinh(-u) = -\sinh u. \quad (65)$$

Substituting $(-x)$ for x in (58) and (61), does not change the value of the square root, that is,

$$\cos(-u) = \cos u, \quad (66)$$

$$\cosh(-u) = \cosh u, \quad (67)$$

Which signifies that $\cos u$ and $\cosh u$ are *even functions*, while $\sin u$ and $\sinh u$ are *odd functions*.

Adding and subtracting (60) and (61), gives

$$e^{\pm u} = \cosh u \pm \sinh u. \quad (68)$$

(b) FUNCTIONS OF IMAGINARY VARIABLES.

188. Substituting, in (56) and (59), $x = -jy$, thus $y = jx$, gives

$$(56.) \quad u = \int \frac{dx}{\sqrt{1-x^2}};$$

$$(59.) \quad u = \int \frac{dx}{\sqrt{1+x^2}};$$

$$x = \sin u;$$

$$x = \sinh u = \frac{\epsilon^u - \epsilon^{-u}}{2};$$

$$\sqrt{1+x^2} = \cos u;$$

$$\sqrt{1+x^2} = \cosh u = \frac{\epsilon^u + \epsilon^{-u}}{2};$$

$$\text{hence, } ju = \int \frac{dy}{\sqrt{1+y^2}},$$

$$\text{hence, } ju = \int \frac{dy}{\sqrt{1-y^2}};$$

$$y = \sinh ju = \frac{\epsilon^{ju} - \epsilon^{-ju}}{2};$$

$$y = \sin ju; \quad . \quad . \quad . \quad (69)$$

$$\sqrt{1+y^2} = \cosh ju = \frac{\epsilon^{ju} + \epsilon^{-ju}}{2};$$

$$\sqrt{1-y^2} = \cos ju; \quad . \quad . \quad . \quad (70)$$

Resubstituting x in both

$$x = \sin u = \frac{\sinh ju}{j} = \frac{\epsilon^{ju} - \epsilon^{-ju}}{2j}; \quad x = \sinh u = \frac{\epsilon^u - \epsilon^{-u}}{2} = \frac{\sin ju}{j}; \quad (71)$$

$$\begin{aligned} \sqrt{1-x^2} = \cos u &= \cosh ju \\ &= \frac{\epsilon^{ju} + \epsilon^{-ju}}{2}; \end{aligned}$$

$$\begin{aligned} \sqrt{1+x^2} = \cosh u &= \frac{\epsilon^u + \epsilon^{-u}}{2} \\ &= \cos ju. \quad . \quad (72) \end{aligned}$$

Adding and subtracting,

$$\epsilon^{\pm ju} = \cos u \pm j \sin u = \cosh ju \pm \sinh ju$$

and

$$\epsilon^{\pm u} = \cosh u \pm \sinh u = \cos ju \mp j \sin ju. \quad . \quad . \quad (73)$$

(c) FUNCTIONS OF COMPLEX VARIABLES

189. It is:

$$\epsilon^{u \pm jv} = \epsilon^u \epsilon^{\pm jv} = \epsilon^u (\cos v \pm j \sin v); \quad . \quad . \quad . \quad (74)$$

$$\left. \begin{aligned} \sin(u \pm jv) &= \sin u \cos jv \pm \cos u \sin jv \\ &= \sin u \cosh v \pm j \cos u \sinh v = \frac{e^v + e^{-v}}{2} \sin u \pm j \frac{e^v - e^{-v}}{2} \cos u; \end{aligned} \right\} \quad (75)$$

$$\left. \begin{aligned} \cos(u \pm jv) &= \cos u \cos jv \mp \sin u \sin jv \\ &= \cos u \cosh v \mp j \sin u \sinh v = \frac{e^v + e^{-v}}{2} \cos u \mp j \frac{e^v - e^{-v}}{2} \sin u; \end{aligned} \right\} \quad (76)$$

$$\left. \begin{aligned} \sinh(u \pm jv) &= \frac{e^{u \pm jv} - e^{-u \mp jv}}{2} = \frac{e^u - e^{-u}}{2} \cos v \pm j \frac{e^u + e^{-u}}{2} \sin v \\ &= \sinh u \cos v \pm j \cosh u \sin v; \end{aligned} \right\} \quad (77)$$

$$\left. \begin{aligned} \cosh(u \pm jv) &= \frac{e^{u \pm jv} + e^{-u \mp jv}}{2} = \frac{e^u + e^{-u}}{2} \cos v \pm j \frac{e^u - e^{-u}}{2} \sin v \\ &= \cosh u \cos v \pm j \sinh u \sin v; \end{aligned} \right\} \quad (78)$$

etc.

(d) RELATIONS.

190. From the preceding equations it thus follows that the three functions, exponential, trigonometric, and hyperbolic, are so related to each other, that any one of them can be expressed by any other one, so that when allowing imaginary and complex imaginary variables, one function is sufficient. As such, naturally, the exponential function would generally be chosen.

Furthermore, it follows, that all functions with imaginary and complex imaginary variables can be reduced to functions of real variables by the use of only two of the three classes of functions. In this case, the exponential and the trigonometric functions would usually be chosen.

Since functions with complex imaginary variables can be resolved into functions with real variables, for their calculation tables of the functions of real variables are sufficient.

The relations, by which any function can be expressed by any other, are calculated from the preceding paragraph, by the following equations:

$$\left. \begin{aligned} \epsilon^{\pm u} &= \cosh u \pm \sinh u = \cos ju \mp j \sin ju; \\ \epsilon^{\pm jv} &= \cos v \pm j \sin v = \cosh jv \pm j \sinh jv; \\ \epsilon^{u \pm jv} &= \epsilon^u (\cos v \pm j \sin v), \end{aligned} \right\} \quad (a)$$

$$\left. \begin{aligned} \sin u &= \frac{\sinh ju}{j} = \frac{\epsilon^{ju} - \epsilon^{-ju}}{2j}; \\ \sin jv &= j \sinh v = j \frac{\epsilon^v - \epsilon^{-v}}{2}; \\ \sin (u \pm jv) &= \sin u \cosh v \pm j \cos u \sinh v \\ &= \frac{\epsilon^v + \epsilon^{-v}}{2} \sin u \pm j \frac{\epsilon^v - \epsilon^{-v}}{2} \cos u; \end{aligned} \right\} \quad (b)$$

$$\left. \begin{aligned} \cos u &= \cosh ju = \frac{\epsilon^{ju} + \epsilon^{-ju}}{2}; \\ \cos jv &= \cosh v = \frac{\epsilon^{jv} + \epsilon^{-jv}}{2}; \\ \cos (u \pm jv) &= \cos u \cosh v \mp j \sin u \sinh v \\ &= \frac{\epsilon^v - \epsilon^{-v}}{2} \cos u \mp j \frac{\epsilon^v - \epsilon^{-v}}{2} \sin u; \end{aligned} \right\} \quad (c)$$

$$\left. \begin{aligned} \sinh u &= \frac{\epsilon^u - \epsilon^{-u}}{2} = \frac{\sin ju}{j}; \\ \sinh jv &= j \sin v = \frac{\epsilon^{jv} - \epsilon^{-jv}}{2}; \\ \sinh (u \pm jv) &= \sinh u \cos v \pm j \cosh u \sin v \\ &= \frac{\epsilon^u - \epsilon^{-u}}{2} \cos v \pm j \frac{\epsilon^u + \epsilon^{-u}}{2} \sin v; \end{aligned} \right\} \quad (d)$$

$$\left. \begin{aligned} \cosh u &= \frac{\epsilon^u + \epsilon^{-u}}{2} = \cos ju; \\ \cosh jv &= \cos v = \frac{\epsilon^{jv} + \epsilon^{-jv}}{2}; \\ \cosh (u \pm jv) &= \cosh u \cos v \pm j \sinh u \sin v \\ &= \frac{\epsilon^u + \epsilon^{-u}}{2} \cos v \pm j \frac{\epsilon^u - \epsilon^{-u}}{2} \sin v. \end{aligned} \right\} \quad (e)$$

And from (b) and (d), respectively (c) and (e), it follows that

$$\left. \begin{aligned} \sinh (u \pm jv) &= j \sin (\pm v - ju) = \pm j \sin (v \pm ju); \\ \cosh (u \pm jv) &= \cos (v \mp ju). \end{aligned} \right\} \quad (f)$$

Tables of the exponential functions and their logarithms, and of the hyperbolic functions with real variables, are given in the following Appendix B.

APPENDIX B.

TWO TABLES OF EXPONENTIAL AND HYPERBOLIC FUNCTIONS.

TABLE I.

$$\varepsilon^{-x}$$

$$\varepsilon=2.7183,$$

$$\log \varepsilon=0.4343.$$

x	$\times 10^{-3}$	$\times 10^{-2}$	$\times 10^{-1}$	$\times 1$
1.0	0.999	0.990	0.905	0.368
1.2		0.988	0.887	0.301
1.4		0.986	0.869	0.247
1.6		0.984	0.852	0.202
1.8		0.982	0.835	0.165
2.0	0.998	0.980	0.819	0.135
2.5		0.975	0.779	0.082
3.0	0.997	0.970	0.741	0.050
3.5		0.966	0.705	0.030
4.0	0.996	0.961	0.670	0.018
4.5		0.956	0.638	0.011
5.0	0.995	0.951	0.607	0.007
6	0.994	0.942	0.549	0.002
7	0.993	0.932	0.497	0.001
8	0.992	0.923	0.449	0.000
9	0.991	0.914	0.407	
10	0.990	0.905	0.368	

TABLE II.

EXPONENTIAL AND HYPERBOLIC FUNCTIONS.

$$e = 2.718282 \sim 2.7183,$$

$$\log e = 0.4342945 \sim 0.4343.$$

$$\cosh x = \frac{1}{2}(e^x + e^{-x}),$$

$$\sinh x = \frac{1}{2}(e^x - e^{-x}).$$

p.p.		
	434	435
0	0	0
0.1	43	43
0.2	87	87
0.3	130	130
0.4	174	174
0.5	217	217
0.6	261	261
0.7	304	304
0.8	347	348
0.9	391	391
1.0	434	435

x	$\log e^x$	$\frac{d \log}{d x} e^x$	$\log e^{-x}$	e^x	e^{-x}	$\cosh x$	$\sinh x$	x
0	0		0	1	1	1	0	0
0.001	0.000434	434	9.999566	1.00100	0.99900	1.00000	0.00100	0.001
0.002	0.000869	435	9.999131	1.00200	0.99800	1.00000	0.00200	0.002
0.003	0.001303	434	9.998697	1.00301	0.99700	1.00000	0.00300	0.003
0.004	0.001737	434	9.998263	1.00401	0.99601	1.00001	0.00400	0.004
0.005	0.002171	434	9.997829	1.00501	0.99501	1.00001	0.00500	0.005
0.006	0.002606	435	9.997394	1.00602	0.99402	1.00002	0.00600	0.006
0.007	0.003040	434	9.996960	1.00702	0.99302	1.00002	0.00700	0.007
0.008	0.003474	434	9.996526	1.00803	0.99203	1.00003	0.00800	0.008
0.009	0.003909	435	9.996091	1.00904	0.99104	1.00004	0.00900	0.009
0.010	0.004343	434	9.995657	1.01005	0.99005	1.00005	0.01000	0.010
0.012	0.005212		9.994788	1.01207	0.98807	1.00007	0.01200	0.012
0.014	0.006080		9.993920	1.01410	0.98610	1.00010	0.01400	0.014
0.016	0.006949		9.993051	1.01613	0.98413	1.00013	0.01600	0.016
0.018	0.007817		9.992183	1.01816	0.98216	1.00016	0.01800	0.018
0.020	0.008686		9.991314	1.02020	0.98020	1.00020	0.02000	0.020
0.025	0.010857		9.989143	1.02531	0.97531	1.00031	0.02500	0.025
0.030	0.013029		9.986971	1.03046	0.97046	1.00046	0.03000	0.030
0.035	0.015200		9.984800	1.03562	0.96562	1.00062	0.03500	0.035
0.040	0.017372		9.982628	1.04081	0.96079	1.00080	0.04001	0.040
0.045	0.019543		9.980457	1.04603	0.95603	1.00102	0.04502	0.045
0.050	0.021715		9.978285	1.05127	0.95123	1.00125	0.05003	0.050
0.06	0.026058		9.973942	1.06184	0.94176	1.00180	0.06004	0.06
0.07	0.030401		9.969599	1.07251	0.93239	1.00245	0.07006	0.07
0.08	0.034744		9.965256	1.08329	0.92312	1.00321	0.08008	0.08
0.09	0.039086		9.960911	1.09417	0.91393	1.00405	0.09011	0.09
0.10	0.043429		9.956571	1.10516	0.90484	1.00500	0.10016	0.10
0.12	0.052115		9.947885	1.12750	0.88692	1.00721	0.12028	0.12
0.14	0.060801		9.939199	1.15027	0.86936	1.00982	0.14046	0.14
0.16	0.069487		9.930513	1.17351	0.85214	1.01283	0.16069	0.16
0.18	0.078173		9.921827	1.19721	0.83527	1.01624	0.18097	0.18
0.20	0.086859		9.913141	1.22140	0.81873	1.02006	0.20131	0.20

$$e^{+0.001} = 1.001000494,$$

$$e^{-0.001} = 0.99900049$$

TABLE II—Continued.

EXPONENTIAL AND HYPERBOLIC FUNCTIONS.

x	$\log e+x$	$\log e^{-x}$	$e+x$	e^{-x}	$\cosh x$	$\sinh x$	x
0.20	0.086859	9.913141	1.22140	0.81873	1.02006	0.20134	0.20
0.25	0.108574	9.891426	1.28403	0.77880	1.03142	0.25261	0.25
0.30	0.130288	9.869712	1.34986	0.74082	1.04534	0.30457	0.30
0.35	0.152003	9.847997	1.41907	0.70469	1.06188	0.35719	0.35
0.40	0.173718	9.826282	1.49183	0.67032	1.08108	0.41076	0.40
0.45	0.195133	9.804567	1.56831	0.63763	1.10297	0.46534	0.45
0.50	0.217147	9.782853	1.64870	0.60653	1.12761	0.52108	0.50
0.6	0.260577	9.739423	1.82212	0.54881	1.19546	0.63666	0.6
0.7	0.304006	9.695994	2.01375	0.49659	1.25517	0.75858	0.7
0.8	0.347436	9.652564	2.22554	0.44933	1.33744	0.88811	0.8
0.9	0.390865	9.609135	2.45960	0.40657	1.43309	1.02657	0.9
1.0	0.434294	9.565706	2.71828	0.36788	1.54308	1.17520	1.0
1.2	0.521153	9.478847	3.32011	0.30119	1.81065	1.50946	1.2
1.4	0.608012	9.391988	4.05520	0.24660	2.15090	1.90430	1.4
1.6	0.694871	9.305129	4.95304	0.20190	2.57745	2.37557	1.6
1.8	0.781730	9.218270	6.04965	0.16530	3.10745	3.44218	1.8
2.0	0.868589	9.131411	7.38906	0.13534	3.76220	3.62686	2.0
2.5	1.085736	8.914264	12.1825	0.082085	6.1323	6.0002	2.5
3.0	1.302883	8.694117	20.0855	0.049797	10.0677	10.0178	3.0
3.5	1.520030	8.479970	33.1154	0.030197	16.5728	16.5426	3.5
4.0	1.737178	8.262822	54.5983	0.018316	27.3083	27.2900	4.0
4.5	1.954325	8.045675	90.0170	0.011109	45.0141	45.0030	4.5
5.0	2.171472	7.828528	148.413	0.006738	74.210	74.203	5.0
6	2.605767	7.394233	103.428	0.002479	201.715	201.713	6
7	3.040061	6.959939	1096.63	0.000912	$=\frac{1}{2}e+x$ for $x > 6$		7
8	3.474356	6.525644	2980.96	0.000335			8
9	3.908650	6.091350	8103.08	0.000123			9
10	4.342945	5.657055	22026.5	0.0000454			10
12	5.211534	4.788166	162755	0.0000061			12
14	6.080123	3.919877	1202610	0.00000083			14
16	6.948712	3.051288	8886120	0.00000011			16
18	7.817301	2.182699	65660000	0.00000002			18
20	8.685890	1.314110	485166000	0.00000000			20

INDEX

A

- Abelian integrals and functions, 305
- Absolute number, 4
 - value of fractional expression, 49
 - of general number, 30
- Accuracy, loss of, 281
 - of approximation estimated, 200
 - of calculation, 279
 - of curve equation, 210
 - of transmission line equations, 208
- Addition, 1
 - of general number, 28
 - and subtraction of trigonometric functions, 102
- Algebra of general number or complex quantity, 25
- Algebraic expression, 294
 - function, 75
- Alternating current and voltage vector, 41
 - functions, 117, 125
 - waves, 117, 125
- Alternations, 117
- Alternator short circuit current, approximated, 195
- Analytical calculation of extrema, 152
 - function, 294
- Angle, see also *Phase angle*.
- Approximation calculation, 280
 - by chain fraction, 208c
- Approximations giving $(1 + s)$ and $(1 - s)$, 201
 - of infinite series, 53
 - methods of, 187
- Arbitrary constants of series, 69, 79
- Area of triangle, 106
- Arrangement of numerical calculations, 275
- Attack, method of, 275

B

- Base of logarithm, 21
- Binomial series with small quantities, 193
 - theorem, infinite series, 59
 - of trigonometric function, 104
- Biquadratic parabola, 219

C

- Calculation, accuracy, 279
 - checking of, 291
 - numerical, 258
 - reliability, 271
- Capacity, 65
- Chain fraction, 208
- Change of curve law, 211, 234
- Characteristics of exponential curves, 228
 - of parabolic and hyperbolic curves, 223
- Charging current maximum of condenser, 176
- Checking calculations, 293a
- Ciphers, number of, in calculations, 282
- Circle defining trigonometric functions, 94
- Coefficients, unknown, of infinite series, 60
- Combination of exponential functions, 231
 - of general numbers, 28
 - of vectors, 29
- Comparison of exponential and hyperbolic curves, 229
- Complementary angles in trigonometric functions, 99
- Complex imaginary quantities, see *General number*.

- Complex, quantity, 17
 - algebra, 27
 - see *General number*.
 - Conjugate numbers, 31
 - Constant, arbitrary of series, 69, 79
 - errors, 186
 - factor with parabolic and hyperbolic curves, 223
 - phenomena, 106
 - terms of curve equation, 211
 - of empirical curves, 234
 - in exponential curves, 230
 - with exponential curves, 229
 - in parabolic and hyperbolic curves, 225
 - Convergency determinations of
 - potential series, 215
 - of series, 57
 - Convergent series, 56
 - Coreless by potential series, 213
 - curve evaluation, 244
 - Cosecant function, 98
 - Cosh function, 305
 - Cosine-amplitude, 299
 - components of wave, 121, 125
 - function, 94
 - series, 82
 - versed function, 98
 - Cotangent function, 94
 - Counting, 1
 - Current change curve evaluation, 241
 - of distorted voltage wave, 169
 - input of induction motor, approximated, 191
 - maximum of alternating transmission circuit, 159
 - Curves, checking calculations, 293*b*
 - empirical, 209
 - law, change, 234
 - rational equation, 210
 - use of, 284
- D
- Data on calculations and curves, 271
 - derived from curve, 285
 - Decimal error, 293*b*
 - Decimals, number of, in calculations, 282
 - in logarithmic tables, 281
 - Definite integrals of trigonometric functions, 103
 - Degrees of accuracy, 279
 - Delta-amplitude, 299
 - Differential equations, 64
 - of electrical engineering, 65, 78, 86
 - of second order, 78
 - Differentiation of trigonometric functions, 103
 - Diophantic equations, 186
 - Distorted electric waves, 108
 - Distortion of wave, 139
 - Divergent series, 56
 - Division, 6
 - of general number, 42
 - with small quantities, 190
 - Double angles in trigonometric functions, 103
 - peaked wave, 255, 260, 266
 - periodicity of elliptic functions, 299
 - scale, 289
- E
- ϵ , 21
 - Efficiency maximum of alternator, 162
 - of impulse turbine, 154
 - of induction generator, 177
 - of transformer, 155, 174
 - Electrical engineering, differential equations, 65, 78, 86
 - Ellipse, length of arc, 301
 - Elliptic integrals and functions, 299
 - Empirical curves, 209
 - evaluation, 233
 - equation of curve, 210
 - Engineering differential equations, 65, 78, 86
 - reports, 290
 - Equilateral hyperbola, 217
 - Errors, constant, 186
 - numerical, 293*b*
 - of observation, 180

- Estimate of accuracy of approximation, 200
 Evaluation of empirical curves, 233
 Even functions, 81, 98, 305
 periodic, 122
 harmonics, 117, 266
 separation, 120, 125, 134
 Evolution, 9
 of general number, 44
 of series, 70
 Exact calculation, 281
 Exciting current of transformer, resolution, 137
 Explicit analytic function, 294
 Exponent, 9
 Exponential curves, 227
 forms of general number, 50
 functions, 52, 297, 304
 with small quantities, 196
 series, 71
 tables, 312, 313, 314
 and trigonometric functions, relation, 83
 Extrapolation on curve, limitation, 210
 Extrema, 147
 analytic determination, 152
 graphical construction of differential function, 170
 graphical determination, 147, 150, 168
 with intermediate variables, 155
 with several variables, 163
 simplification of function, 157
- F
- Factor, constant, with parabolic and hyperbolic curves, 223
 Fan motor torque by potential series, 215
 Fifth harmonic, 261, 264
 Flat top wave, 255, 260, 265, 268
 zero waves, 255, 258, 261, 265
 Fourier series, see *Trigonometric series*.
 Fraction, 8
 as series, 52
 chain-, 208
- Fractional exponents, 11, 44
 expressions of general number, 49
 Full scale, 289
 Functions, theory of, 294
- G
- Gamma function, 304
 General number, 1, 16
 algebra, 25
 engineering reports, 291
 exponential forms, 50
 reduction, 48
 reports on engineering matters, 292
 Geometric scale of curve plotting, 288
 Graphical determination of extrema, 147, 150, 168
- H
- Half angles in trigonometric functions, 103
 Half waves, 117
 Half scale, 289
 Harmonics, even, 117
 odd, 117
 of trigonometric series, 114
 two, in wave, 255
 High harmonics in wave shape, 255, 269
 Hunting of synchronous machines, 257
 Hyperbola, arc of, 61
 equilateral, 217
 Hyperbolic curves, 216
 functions, 294
 curve, shape, 232
 integrals and functions, 298
 tables, 313, 314
 Hyperelliptic integrals and functions, 301
 Hysteresis curve of silicon steel, investigation of, 248
- I
- Imaginary number, 26
 quantity, see *Quadrature number*.

Incommensurable waves, 257
 Indeterminate coefficients, method, 71
 Indeterminate coefficients of infinite series, 60
 Individuals, 8
 Inductance, 65
 Infinite series, 52
 values of curves, 211
 of empirical curves, 233
 Inflection points of curves, 153
 Impedance vector, 41
 Implicit analytic function, 294
 Integral function, 295
 Integration constant of series, 69, 79
 of differential equation, 65
 by infinite series, 60
 of trigonometric functions, 103
 Intelligibility of calculations, 283
 Intercepts, defining tangent and co-tangent functions, 94
 Involution, 9
 of general numbers, 44
 Irrational numbers, 11
 Irrationality of representation by potential series, 213

J

J, 14

L

Least squares, method of, 179, 186
 Limitation of mathematical representation, 40
 of method of least squares, 186
 of potential series, 216
 Limiting value of infinite series, 54
 Linear number, 33
 see *Positive and Negative number*.
 Line calculation, 276
 equations, approximated, 204
 Logarithm of exponential curve, 229
 as infinite series, 63
 of parabolic and hyperbolic curves, 225
 with small quantities, 197

Logarithmation, 20
 of general numbers, 51
 Logarithmic curves, 227
 functions, 297
 paper, 233, 287
 scale, 288
 tables, number of decimals in, 281
 Loss of curve induction motor, 183

M

Magnetic characteristic on semi-logarithmic paper, 288
 Magnetite arc, volt-ampere characteristic, 239
 characteristic, evaluation, 246
 Magnitude of effect, determination, 280
 Maximum, see *Extremum*.
 Maxima, 147
 McLaurin's series with small quantities, 198
 Mechanism of calculating empirical curves, 237
 Methods of calculation, 275
 of intermediate coefficients, 71
 of least squares, 179, 186
 of attack, 275
 Minima, 147
 Minimum, see *Extremum*.
 Multiple frequencies of waves, 274
 Multiplicand, 39
 Multiplication, 6
 of general numbers, 39
 with small quantities, 188
 of trigonometric functions, 102
 Multiplier, 39

N

Negative angles in trigonometric functions, 98
 exponents, 11
 number, 4
 Nodes in wave shape, 256, 270
 Non-periodic curves, 212
 Nozzle efficiency, maximum, 150
 Number, general, 1

Numerical calculations, 275
 values of trigonometric functions,
 101

O

Observation, errors, 180
 Octave as logarithmic scale, 288
 Odd functions, 81, 98, 305
 period c , 122
 harmonics in symmetrical wave,
 117
 separation, 120, 125, 134
 Omissions in calculations, 293b
 Operator, 40
 Order of small quantity, 188
 Oscillating functions, 92
 Output, see *Power*.

P

π and $\frac{\pi}{2}$ added and subtracted in
 trigonometric function, 100
 approximated by chain fraction,
 208c
 Pairs of high harmonics, 270
 Parabola, common, 218
 Parabolic curves, 216
 Parallelogram law of general num-
 bers, 28
 of vectors, 29
 Peaked wave, 255, 258, 261, 264
 Pendulum motion, 301
 Percentage change of parabolic and
 hyperbolic curves, 223
 Periodic curves, 254
 decimal fraction, 12
 phenomena, 106
 Periodicity, double, of elliptic func-
 tions, 299
 of trigonometric functions, 96
 Permeability maximum, 148, 170
 Phase angle of fractional expression,
 49
 of general number, 28
 Plain number, 32
 see *General number*.

Plotting of curves, 212
 proper and improper, 286
 of empirical curve, 234
 Polar co-ordinates of general num-
 ber, 25, 27
 expression of general number, 25,
 27, 38, 43, 44, 48
 Polyphase relation, reducing trigo-
 nometric series, 134
 of trigonometric functions, 104
 system of points or vectors, 46
 Positive number, 4
 Potential series, 52, 212
 Power factor maximum of induction
 motor, 149
 maximum of alternating trans-
 mission circuit, 158
 of generator, 161
 of shunted resistance, 155
 of storage battery, 172
 of transformer, 173
 of transmission line, 165
 not vector product, 42
 of shunt motor, approximated, 189
 with small quantities, 194
 Probability calculation, 181
 Product series, 303
 of trigonometric functions, 102
 Projection, defining cosine function,
 94
 Projector, defining sine function, 94

Q

Quadrants, sign of trigonometric
 functions, 96
 Quadrature numbers, 13
 Quarter scale, 289
 Quaternions, 22

R

Radius vector of general number, 28
 Range of convergency of series, 56
 Ratio of variation, 226
 Rational equation of curve, 210
 function, 295
 Rationality of potential series, 214

- Real number, 26
 - Rectangular co-ordinates of general number, 25
 - Reduction to absolute values, 48
 - Relations of hyperbolic trigonometric and exponential functions, 309
 - Relativeness of small quantities, 188
 - Reliability of numerical calculations, 293
 - Reports, engineering, 290
 - Resistance, 65
 - Resolution of vectors, 29
 - Reversal by negative unit, 14
 - double, at zero of wave, 258, 261
 - Reverse function, 294
 - Right triangle defining trigonometric functions, 94
 - Ripples in wave, 45
 - by high harmonics, 270
 - Roots of general numbers, 45
 - expressed by periodic chain fraction, 208*e*
 - with small quantities, 194
 - of unit, 18, 19, 46
 - Rotation by negative unit, 14
 - by quadrature unit, 14
- S
- Saddle point, 165
 - Saw-tooth wave, 246, 255, 258, 260, 265
 - Scalar, 26, 28, 30
 - Scale in curve plotting, proper and improper, 212, 286
 - full, double, half, etc., 287
 - Scientific engineering records, 291
 - Secant function, 98
 - Second harmonic, effect of, 266
 - Secondary effects, 210
 - phenomena, 234
 - Semi-logarithmic paper, 287
 - Series, exponential, 71
 - infinite, 52
 - trigonometric, 106
 - Seventh harmonic, 262
 - Shape of curves, 212
 - proper in plotting, 286
 - of exponential curve, 227, 230
 - of function, by curve, 284
 - of hyperbolic functions, 232
 - of parabolic and hyperbolic curves, 217
 - Sharp zero wave, 255, 260, 265
 - Short circuit current of alternator, approximated, 195
 - Sign error, 293*c*
 - of trigonometric functions, 95
 - Silicon steel, investigation of hysteresis curve, 248
 - Simplification by approximation, 187
 - Sine-amplitude, 199
 - component of wave, 121, 125
 - function, 94
 - series, 82
 - versus function, 98
 - Sine function, 305
 - Slide rule accuracy, 281
 - Small quantities, approximation, 187
 - Special functions, 302
 - Squares, least, method of, 179, 186
 - Steam path of turbine, 33
 - Subtraction, 1
 - of general number, 28
 - of trigonometric functions, 102
 - Summation series, 303
 - Superposition of high harmonics, 273
 - Supplementary angles in trigonometric functions, 99
 - Surging of synchronous machines, 301
 - Symmetrical curve maximum, 150
 - periodic function, 117
 - wave, 117
- T
- Tabular form of calculation, 275
 - Tangent function, 94
 - Taylor's series with small quantities, 199
 - Temperature wave, 131
 - Temporary use of potential series, 216

- Terminal conditions of problem, 69
Terms, constant, of empirical curves, 234
 in exponential curve, 229
 with exponential curve, 229
 in parabolic and hyperbolic curves, 225
 of infinite series, 53
Theorem, binomial, infinite series, 59
Thermomotive force wave, 133
Theta functions, 300
Third harmonic, 136, 255
Top, peaked or flat, of wave, 255
Torque of fan motor by potential series, 215
Transient current curve, evaluation, 241
 phenomena, 106
Transmission equations, approximated, 204
 line calculation, 275
Treble peak of wave, 262
Triangle, defining trigonometric functions, 94
 trigonometric relations, 106
Trigonometrical and exponential functions, relations, 83
 functions, 94, 304
 series, 82
 with small quantity, 198
 integrals and functions, 298
 series, 106
 calculation, 114, 116, 139
Triple harmonic, separation, 136
 peaked wave, 255
 scale, 289
Tungsten filament, volt-ampere characteristic, 235
Turbine, steam path, 33

U
Unequal height and length of half waves, 268
Univalent functions, 106
Unsymmetric curve maximum, 151
 wave, 138

V
Values of trigonometric functions, 101
Variation, ratio of, 226
Vector analysis, 32
 multiplication, 39
 quantity, 32
 see General number.
 representation by general number, 29
Velocity diagram of turbine steam path, 34
 functions of electric field, 304
Versed sine and cosine functions, 98
Volt-ampere characteristic of magnetite arc, 239
 of tungsten filament, 235

Z
Zero values of curve, 211
 of empirical curves, 233
 of waves, 255







**PLEASE DO NOT REMOVE
CARDS OR SLIPS FROM THIS POCKET**

UNIVERSITY OF TORONTO LIBRARY

P&A Sci.

